

Sparse Autoencoders for Hypothesis Generation

Rajiv Movva^{*1} Kenny Peng^{*23} Nikhil Garg³ Jon Kleinberg² Emma Pierson¹

Abstract

We describe HYPOTHESAES, a general method to hypothesize interpretable relationships between text data (e.g., headlines) and a target variable (e.g., clicks). HYPOTHESAES has three steps: (1) train a sparse autoencoder on text embeddings to produce interpretable features describing the data distribution, (2) select features that predict the target variable, and (3) generate a natural language interpretation of each feature (e.g., *mentions being surprised or shocked*) using an LLM. Each interpretation serves as a hypothesis about what predicts the target variable. Compared to baselines, our method better identifies reference hypotheses on synthetic datasets (at least +0.06 in F1) and produces more predictive hypotheses on real datasets (~twice as many significant findings), despite requiring 1-2 orders of magnitude less compute than recent LLM-based methods. HYPOTHESAES also produces novel discoveries on two well-studied tasks: explaining partisan differences in Congressional speeches and identifying drivers of engagement with online headlines.

1. Introduction

Large language models (LLMs) show promise as a tool for *hypothesis generation*. Discovering relationships between text data and a target variable is an important and fundamental task with diverse applications in economics, political science, sociology, medicine, and business (Grimmer, 2010; Rathje et al., 2021; Sun et al., 2022; Gentzkow & Shapiro, 2010; Monroe et al., 2009; Nelson, 2020; Ranard et al., 2016; Ting et al., 2017). What features of a restaurant review predict a low rating? What features of a social media post predict whether it will go viral? What features of a patient’s clinical notes predict if they will develop cancer?

Automated approaches to answering these questions have

the potential to significantly expand and accelerate scientific discovery (Ludwig & Mullainathan, 2023). Indeed, multiple lines of work—spanning decades of research—have sought to extract text features that can predict a target variable (Blei et al., 2003; Monroe et al., 2009). Recent LLM-based hypothesis generation methods are especially exciting since they operate directly at the level of human-interpretable natural language concepts (Zhou et al., 2024; Batista & Ross, 2024; Ludan et al., 2024; Zhong et al., 2024). Concretely, hypothesis generation methods take a dataset of texts linked to a target variable (e.g., headlines and their engagement level) and hypothesize natural language concepts that predict the target variable (e.g., *mentions being surprised or shocked* or *asks a question to the reader*).

However, basic hurdles impede the full usage of LLMs for hypothesis generation. Consider two natural approaches. The first prompts an LLM with positive and negative examples and asks it to hypothesize concepts that distinguish them (Zhou et al., 2024; Batista & Ross, 2024; Ludan et al., 2024). However, due to context window and reasoning limitations, we can only give LLMs a few training examples at a time, making this hard to scale. A second approach aims to interpret a model fine-tuned to predict the target variable (Zhong et al., 2024), leveraging the full training dataset. However, neurons are often hard to interpret (Elhage et al., 2022). In summary, efficient and scalable learning of statistical relationships requires numerical representations of text (e.g., *neurons*), but interpreting neurons is hard.

Our first contribution is a theoretical framework that connects model interpretation and hypothesis generation, clarifying the challenge above. Proposition 3.1 gives a “triangle inequality” for hypothesis generation: if a neuron is predictive of the target variable, then a sufficiently high-fidelity interpretation of the neuron is also predictive. This result directly motivates a general hypothesis generation procedure:

1. (Feature generation) Learn interpretable neurons (i.e., that tend to fire in the presence of human concepts).
2. (Feature selection) Select neurons that predict the target variable.
3. (Feature interpretation) Generate high-fidelity natural language interpretations of these neurons.

^{*}Equal contribution ¹UC Berkeley ²Cornell University ³Cornell Tech. Correspondence to: Rajiv Movva <rmovva@berkeley.edu>, Kenny Peng <kennypeng@cs.cornell.edu>.

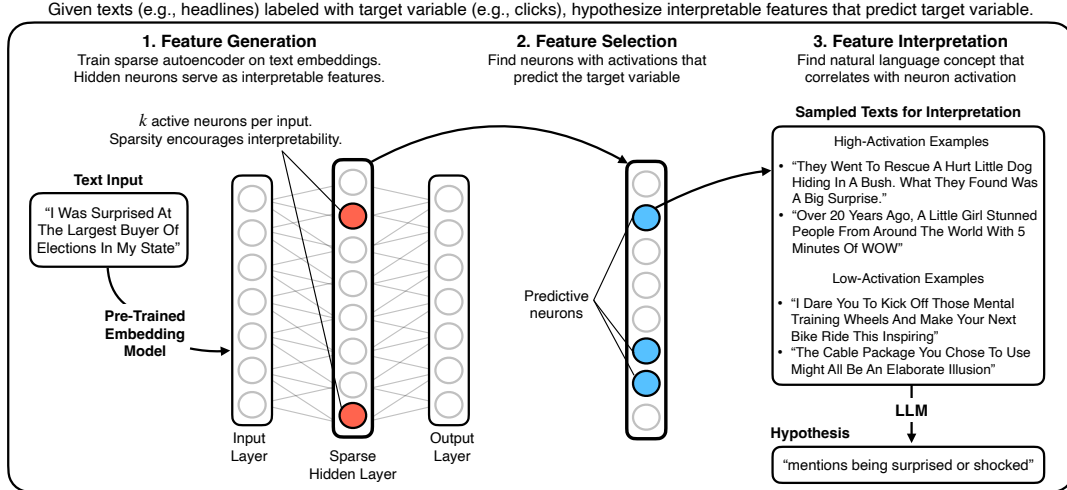


Figure 1. HYPOTHESAES is a general method that outputs natural language hypotheses from text datasets.

Our second contribution is a method, HYPOTHESAES, that successfully implements this procedure. A primary challenge is learning the interpretable neurons in step 1—after all, a long line of work has established the challenge of interpreting neural networks (Olah et al., 2017; Kim et al., 2018; Elhage et al., 2022). To overcome this, we train a sparse autoencoder (SAE; Makhzani & Frey (2014)) on pre-trained text embeddings from our dataset. Neurons in the SAE hidden layer are highly interpretable, while retaining the explanatory power of embeddings (Cunningham et al., 2023; Bricken et al., 2023; Templeton et al., 2024; Gao et al., 2024; O’Neill et al., 2024). Intuitively, this architecture leverages the fact that natural language is sparse; only a small fraction of all human concepts are expressed at a time. In step 2, we select neurons that are predictive of the target variable (e.g., using Lasso). In step 3, we automatically generate high-fidelity interpretations of these predictive neurons by prompting an LLM with texts that activate the neuron, building on recent work in autointerpretability (Bills et al., 2023). These interpretations serve as hypotheses.

We evaluate HYPOTHESAES against recent state-of-the-art methods. We consider three synthetic tasks, as well as three real-world tasks of practical interest: hypothesizing the relationship between headlines and engagement, speech text and political party, and review text and rating. Example HYPOTHESAES hypotheses are given below:

Task	Example Hypotheses (Abbreviated)
Headline Engagement	(+) <i>mentions being surprised or shocked</i> (-) <i>asks a question directly to the reader</i>
Speaker Party	(Rep.) <i>discusses illegal immigration</i> (Dem.) <i>criticizes tax breaks for the wealthy</i>
Restaurant Rating	(+) <i>mentions plans to return to the restaurant</i> (-) <i>mentions issues related to food safety</i>

HYPOTHESAES produces many more significant hypotheses than three baseline methods. On three real-world tasks, 45/60 hypotheses generated by our method are significant, compared to at most 24 for the baselines. On two datasets, we identify new hypotheses which previous detailed analyses have not uncovered. Intuitively, by decoupling feature generation from the prediction task, HYPOTHESAES is able to explore a wider range of hypotheses.

HYPOTHESAES is also more than $10\times$ faster and cheaper than recent LLM baselines, with similar runtime and cost to BERTopic, a standard embedding-based topic model. Filtering hypotheses based on the predictiveness of neurons is significantly cheaper than directly evaluating natural language concepts, since all neuron activations of an input can be computed from a single forward pass of an SAE (in comparison to one LLM call per natural language feature).

2. Related Work

2.1. Hypothesis Generation

Classical text analysis methods, such as comparing n -gram frequencies between groups (Fightin’ Words; Monroe et al. (2009)) or topic modeling (LDA, Blei et al. (2003); BERTopic, Grootendorst (2022)) remain standard tools to discover relationships between text and target variables (Grimmer et al., 2022). A challenge with these methods is that their outputs—lists of words or documents—are not immediately interpretable by humans (Chang et al., 2009). As a result, applying these methods for scientific discovery requires sufficient domain expertise to prespecify hypotheses and assess whether the data support them (Demszky et al., 2019; Gentzkow et al., 2016; Sun et al., 2022).

Recently, several LLM-based methods aim to automate in-

interpretable concept discovery from text datasets. Grootendorst (2024) proposes using LLMs to assign topics an interpretable label, while Pham et al. (2024) use LLMs for end-to-end topic modeling. Ludan et al. (2024) directly prompt LLMs with a task description to propose interpretable concepts. Sun et al. (2024) and Zhong et al. (2024) also use LLMs to propose concepts, guided by neurons from a supervised model. Zhou et al. (2024) and Batista & Ross (2024) focus specifically on hypothesis generation: LLMs propose natural language descriptions of predictive relationships, which can then be tested for generalization.

Unlike several of these methods, HYPOTHESAES decouples feature generation (unsupervised, using the SAE), selection (supervised), and interpretation (the only LLM-dependent step). Comparing to baselines from each class of methods (§5.3), we illustrate quantitative and qualitative improvements over classical and fully-LLM-driven baselines.

2.2. Sparse Autoencoders

A recent line of work demonstrates that sparse autoencoders (Makhzani & Frey, 2014) are an effective architecture to learn interpretable features from intermediate layers of large language models (Cunningham et al., 2023; Bricken et al., 2023; Templeton et al., 2024; Bills et al., 2023). This literature is primarily motivated by mechanistic interpretability: since the neurons in the hidden layer of sparse autoencoders correspond to interpretable features, they can be strategically altered to steer the behavior of language models. In the present work, we instead apply sparse autoencoders to learn interpretable text features for the purpose of hypothesizing interpretable relationships between text data and a target variable. We follow O’Neill et al. (2024), who show that SAE features can guide semantic search applications, in training a sparse autoencoder directly on text embeddings.

3. Theoretical Framework

In this section, we derive a theoretical result that shows how a language model with highly interpretable neurons can be used for hypothesis generation, as well as where performance loss arises in such a procedure.

Let Z be an indicator for whether a neuron fires for a given text input (i.e., whether the activation exceeds a threshold). Let $\hat{Z}$ be an indicator for whether the text contains a specified natural language concept. Let Y be the target variable. Define the predictiveness of $\hat{Z}$ as the separation score

$$S(\hat{Z}) := |\mathbb{E}[Y|\hat{Z} = 1] - \mathbb{E}[Y|\hat{Z} = 0]|.$$

Define $S(Z)$ analogously.¹ We measure the fidelity of the

¹Separation score is one way we evaluate hypotheses empirically. We suspect that analogs of Proposition 3.1 hold for alternative measures of predictiveness, like explained variance.

interpretation $\hat{Z}$ by

$$\delta(\hat{Z}, Z) := \frac{1 - \min\{\text{recall}, \text{precision}\}}{\min\{\Pr[\hat{Z} = 0], \Pr[Z = 0]\}}, \quad (1)$$

where precision is how often the neuron fires when the concept is present and recall is how often the concept is present when the neuron fires. Lower δ implies higher fidelity. High recall and precision means the interpretation approximates the neuron’s behavior, and is not overly narrow or generic. Empirically, the denominator $\min\{\Pr[\hat{Z} = 0], \Pr[Z = 0]\}$ is close to 1 (neurons and concepts activate infrequently), so maximizing fidelity is approximately maximizing the minimum of recall and precision.

Proposition 3.1 (A Triangle Inequality for Hypothesis Generation). *Suppose Y is supported on $[0, 1]$. Then*

$$|S(\hat{Z}) - S(Z)| \leq \delta(\hat{Z}, Z). \quad (2)$$

Intuitively, the result shows that if $\hat{Z}$ is a high-fidelity predictor of Z , then it will also have a similar separation score. A full proof is given in Appendix E.²

Notice that what we ultimately care about is $S(\hat{Z})$ (how predictive our actual natural language hypothesis is), but our procedure focuses primarily on generating, selecting, and then interpreting the neuron Z . Proposition 3.1 shows that as long as we can generate a sufficiently high-fidelity interpretation $\hat{Z}$ of each neuron Z , then we can focus on finding predictive neurons as a way to ultimately find natural language hypotheses. In particular, the result shows that

$$S(\hat{Z}) = S(Z) + (S(\hat{Z}) - S(Z)) \quad (3)$$

$$\geq \underbrace{S(Z)}_{\text{neuron predictiveness}} - \underbrace{\delta(\hat{Z}, Z)}_{\text{interpretation fidelity}}. \quad (4)$$

Equation (4) shows that to find a natural language concept such that $S(\hat{Z})$ is large, it suffices to first identify Z such that $S(Z)$ is high (i.e., find a predictive neuron) and then $\hat{Z}$ such that $\delta(\hat{Z}, Z)$ is small (i.e., find a high-fidelity interpretation).

This directly motivates the actual procedure employed in HYPOTHESAES. First, we learn interpretable neurons Z using an SAE. Second, we identify which neurons Z are predictive. Third, we generate high-fidelity interpretations $\hat{Z}$ of these neurons with an LLM.

Equation (4) also gives us a lens by which to analyze HYPOTHESAES empirically. The result decomposes the performance of hypotheses $\hat{Z}$ into two parts: the predictiveness of

²There are two main observations that enable us to show the result. First, we focus on $|\mathbb{E}[Y|\hat{Z} = 1] - \mathbb{E}[Y|\hat{Z} = 0]|$, and show that if the false negative rate (FNR) is greater than the false discovery rate (FDR), this value is bounded by the FNR; otherwise, it is bounded by the FDR. Second, we show roughly that as long as $\Pr[\hat{Z} = 1], \Pr[Z = 1] < \Pr[\hat{Z} = 0], \Pr[Z = 0]$, then $|\mathbb{E}[Y|\hat{Z} = 0] - \mathbb{E}[Y|\hat{Z} = 1]| < |\mathbb{E}[Y|\hat{Z} = 1] - \mathbb{E}[Y|\hat{Z} = 0]|$.

the neurons Z and the fidelity of the interpretations $\hat{Z}$. For example, if the SAE neurons are similarly predictive to the text embeddings they encode, then most performance loss (compared to predicting directly from embeddings) comes from low-fidelity neuron interpretation. Imperfect neuron interpretations can be due to either neurons being fundamentally uninterpretable (there doesn't *exist* a high-fidelity interpretation), or due to difficulties in actually *finding* the interpretation. These challenges arise in steps 1 (training the SAE) and step 3 (generating the interpretation) respectively. We perform such an empirical analysis in Appendix C.3.

4. Methods

In the previous section, we established theoretically that we can generate hypotheses by learning interpretable neurons, identifying which neurons are predictive, and then generating high-fidelity interpretations of the neurons. We now describe how we can accomplish these steps in practice.

We now describe the details of HYPOTHESAES. Given a dataset of training examples $\{(x_i, y_i)\}_{i \in [N]}$, where x_i is the input text and y_i is the target variable annotation, HYPOTHESAES outputs H natural language concepts that serve as hypotheses. The goal is for these natural language concepts to predict the target variable.

4.1. Feature Generation: Training Sparse Autoencoders.

Let e_i denote a D -dimensional text embedding of x_i . We train a k -sparse autoencoder (Makhzani & Frey, 2014). k -sparse autoencoders have successfully generated interpretable features both when applied to intermediate layers of language models, as well as text embeddings (Gao et al., 2024; O'Neill et al., 2024). The SAE encodes a text embedding e_i as follows (we use OpenAI's text-embedding-3-small, where $D = 1536$, in our experiments):

$$z_i = \text{ReLU}(\text{TopK}(W_{\text{enc}}(e_i - b_{\text{pre}}) + b_{\text{enc}})) \quad (5)$$

$$\hat{e}_i = W_{\text{dec}} z_i + b_{\text{dec}}, \quad (6)$$

where $b_{\text{pre}} \in \mathbb{R}^D$, $W_{\text{enc}} \in \mathbb{R}^{M \times D}$, $b_{\text{enc}} \in \mathbb{R}^M$, $W_{\text{dec}} \in \mathbb{R}^{D \times M}$, $b_{\text{dec}} \in \mathbb{R}^D$, and TopK sets all activations except the top k to zero. In a basic k -sparse autoencoder, the loss on one input is

$$\mathcal{L} = \|e_i - \hat{e}_i\|_2^2. \quad (7)$$

We follow Gao et al. (2024); O'Neill et al. (2024) in further adding an auxiliary loss to avoid "dead latents." We describe details and hyperparameters in Appendix B.

Empirically, the dimensions of z_i are often highly interpretable, corresponding to human concepts. M sets the total number of concepts learned across the entire dataset, and k sets the number of concepts which can be used to reconstruct each instance. The output of this first step is an $N \times M$ activation matrix, Z_{SAE} .

4.2. Feature Selection: Target Prediction with SAE Neurons.

To select a predictive subset of SAE neurons, we fit an L_1 -regularized linear or logistic regression predicting the target variable Y from the activation matrix Z_{SAE} (Tibshirani, 1996). Formally, for regression³, we optimize

$$\min_{\beta} \mathcal{L}(\beta; \lambda) = \frac{1}{N} \|\mathbf{y} - Z_{\text{SAE}}\beta\|_2^2 + \lambda \|\beta\|_1,$$

where $\mathbf{y}$ is a length- N vector, Z_{SAE} is an $N \times M$ matrix, and β is a length- M vector of feature coefficients for each SAE neuron. The L_1 penalty produces sparse coefficient vectors β where some coefficients are exactly zero (indicating a dropped feature). To generate H hypotheses, we perform binary search to identify a value of λ which produces exactly H nonzero coefficients.

4.3. Feature Interpretation: Labeling Neurons with LLMs.

In this step, we generate high-fidelity interpretations (i.e., natural language concepts) of the subset of predictive neurons. To generate interpretations of a neuron, we prompt an LLM (GPT-4o in our experiments) with example texts from a range of neuron activations and instruct it to identify a natural language concept that is present in high-activation texts and absent in low-activation texts. We then generate multiple interpretations by rerunning this procedure several times at temperature 0.7, and choose the best interpretation according to a measure of fidelity we describe below.

We test a variety of approaches for generating interpretations. In our final procedure, we sample high- and low-activating texts from the neuron's positive activation distribution, with the sampling bins given in Appendix B.2. (We restrict our sample to texts that activate the neuron to ensure some examples contain the concept.) Treating these as true positives and true negatives, respectively, we prompt an LLM to generate a natural language concept that distinguishes 10 samples of each class. We then define an interpretation's fidelity as its F1 score in this prediction task, using an LLM (GPT-4o-mini in our experiments) to annotate 100 positive and 100 negative samples for the presence of the concept.

This approach produced predictive interpretations when compared to other prompting and evaluation schemes, though some other methods performed similarly (see Appendix C for details). Importantly, we found experimentally that higher-fidelity interpretations according to our measure have better predictive performance (Appendix C.1), justifying our approach to selecting interpretations.

³For classification we analogously use an L_1 -regularized loss: $\mathcal{L}(\beta; \lambda) = \frac{1}{N} \text{BCE}(\mathbf{y}, \sigma(Z_{\text{SAE}}\beta)) + \lambda \|\beta\|_1$, where BCE is the binary cross-entropy loss and $\sigma(\cdot)$ is the sigmoid function.

5. Experiments

5.1. Datasets

We evaluate on synthetic and real-world datasets, described below. Appendix A provides more details on all datasets.

Synthetic datasets: recovering known, reference hypotheses. Our synthetic evaluation is motivated by real-world settings in which there are multiple disjoint hypotheses we would like to discover. We therefore use two datasets from prior work on interpretable clustering (Pham et al., 2024; Zhong et al., 2024): **WIKI** and **BILLS**. We construct a target variable that is positive if a text belongs to one of a pre-specified list of reference categories; these are the ground-truth “hypotheses” to recover. Both datasets are labeled by humans with granular hierarchical topics, such as *Media and Drama: Television: The Simpsons* in WIKI and *Macroeconomics: Tax Code* in BILLS.

We generate ground-truth labels based on the most frequent granular topics: articles in any of the top-5 topics are labeled ‘1’, and all others are labeled ‘0’. This produces 5 reference hypotheses, e.g., *the bill is about Macroeconomics: Tax Code* is one of them. For WIKI, we also include a more challenging variation of recovering the top-15 most frequent topics. WIKI contains 11,979 total articles (15% positive for WIKI-5; 35% for WIKI-15), and BILLS contains 21,149 (24% positive). For both datasets, we reserve 2,000 items for validation (i.e., SAE hyperparameter selection) and 2,000 heldout items to evaluate hypotheses; we use the remaining items for SAE training and feature selection.

Real-world datasets: Generating hypotheses that predict a target variable. We apply HYPOTHESAES to three real-world datasets which have been studied in prior work:

- **HEADLINES** (Matias et al., 2021): Which features of digital news headlines predict user engagement? Each instance is a pair of differently-phrased headlines for the same article; users on Upworthy.com were randomly shown one, and the task is to identify which headline received higher click-rate. We reserve a large heldout set to improve statistical power: our split sizes are 8.8K training, 1K validation, and 4.4K heldout.
- **YELP** (Yelp, 2024): Which features of Yelp restaurant reviews predict users’ 1-5 star ratings? We use 200K reviews for training, 10K for validation, and 10K for heldout eval.
- **CONGRESS** (Gentzkow & Shapiro, 2010): Which features of U.S. congressional speeches predict party affiliation? Speeches are from the 109th Congress (2005–07) with binary labels (Rep. or Dem.). Our split sizes are 114K training, 16K validation, and 12K heldout.

On each of these tasks, our goal is to identify interpretable natural language features that predict the target variable. The former two are studied in prior work on hypothesis generation (Zhou et al., 2024; Batista & Ross, 2024; Ludan et al., 2024); the latter is a longstanding application in computational social science (Grimmer et al., 2021).

5.2. Metrics

Synthetic datasets. On the synthetic datasets, we evaluate how well we recover the reference hypotheses. We run HYPOTHESAES by setting H to the number of references (e.g., 5 for WIKI-5). We use GPT-4o-mini to annotate each inferred hypothesis on the heldout set, producing an $N_{\text{heldout}} \times H$ binary matrix of annotations. Following Zhong et al. (2024), we compute the optimal matching between reference and inferred hypotheses via the Hungarian algorithm on the $H \times H$ correlation matrix of reference vs. inferred hypothesis annotations, and we report three metrics:

- **F1 Similarity:** For each matched (reference, inferred) pair, we compute the F1-score between the presence of the ground-truth reference topic and the inferred hypothesis annotations. We report the mean across the H pairs.
- **Surface Similarity:** For each matched pair, we prompt GPT-4o to assess whether the hypotheses are the same, related, or distinct, with corresponding scores of 1.0, 0.5, 0.0. To improve stability, we sample 5 outputs at temperature 0.7 and average them. We report the mean value of this score across the H pairs. (Table 4 shows examples; Figure 5 provides the prompt.)
- **Overall AUC:** Using annotations from the H inferred hypotheses, we fit a multivariate logit and compute AUC to evaluate predictive performance. (Recall that the ground-truth label is a logical OR of the reference topics.)

The first and second metrics are taken directly from Zhong et al. (2024); they assess whether the inferred hypotheses qualitatively and quantitatively match the ground-truth. The third metric assesses overall prediction quality. Some hypotheses may not match a reference but still be predictive, so this metric rewards any interpretable hypotheses which contribute predictive value, even if they are not optimal.

Real datasets. On the real datasets, we generate 20 hypotheses per method, and assess the hypothesis sets for two qualities: (1) *breadth*: how many distinct, predictive hypotheses were identified? (2) *predictiveness*: how well do the hypotheses collectively explain the label? We again compute an $N_{\text{heldout}} \times H$ annotation matrix for each method’s hypotheses, and measure these constructs as follows:

- **Breadth:** We fit a multivariate logit (for binary labels) or OLS (continuous) regressing the target variable against

Source	Overall AUC	# Sig	Significant Hypotheses: Green: ↑ clicks ; Red: ↓ clicks	Sep.	AUC
HYPOTHESAES	0.693	15	mentions photography, images, or visual representations	0.26	0.51
			mentions situations or behaviors that evoke discomfort or embarrassment	0.16	0.58
			includes a reference to being surprised, shocked, or unprepared	0.16	0.56
			focuses on a female subject and her personal journey	0.14	0.52
			mentions or heavily emphasizes a video or video-related content	0.13	0.53
			mentions a surprising or unexpected action or outcome	0.13	0.58
			focuses on a male protagonist or subject	0.11	0.53
			contains phrases about looking or seeing something	0.09	0.52
			mentions uniqueness or exclusivity using superlative or comparative language	0.09	0.53
			mentions a response to or confrontation with offensive behavior	0.09	0.53
			mentions politicians, government actions, or political topics	-0.10	0.52
			asks a question directly to the reader	-0.12	0.53
			addresses collective human responsibility or action	-0.13	0.56
			mentions environmental issues or ecological consequences	-0.15	0.51
			mentions actions or initiatives that positively impact a community	-0.16	0.53
BERTOPIC	0.619	1	headlines emphasize emotional reactions to short, impactful videos	0.15	0.59
NLPARAM	0.660	5	has an element of surprise in its structure	0.14	0.57
			frames the content within a gender-specific context	0.11	0.53
			includes an inspiring transformation	-0.08	0.51
			uses a play on words or puns	-0.09	0.54
HYPOGENIC	0.646	3	poses a rhetorical question	-0.12	0.55
			implies a personal story or anecdote that invites curiosity	0.15	0.56
			uses strong emotional language or evokes curiosity	0.12	0.58
			poses a rhetorical question that may disengage the reader	-0.11	0.55

Table 1. Hypotheses on HEADLINES with $p < 0.05$ after Bonferroni correction in a multivariate regression. ‘Sep.’ is the separation score (§3): $E[Y|\hat{Z} = 1] - E[Y|\hat{Z} = 0]$, where positive scores predict higher click-rates. The ‘AUC’ column provides AUCs of individual hypotheses. ‘Overall AUC’ uses all 20 hypotheses, including insignificant ones. Hypotheses for all methods are abbreviated for space.

all of the hypothesis annotations, and report the count of hypotheses which have a significant, nonzero coefficient⁴. This metric benefits distinct hypotheses and penalizes redundancy. Breadth is important since the utility of hypotheses is often determined by criteria beyond predictiveness. For example, a hypothesis may be more useful if it is connected to past theories or suggests a new direction for investigation. More breadth increases the likelihood that a hypothesis satisfying these criteria is produced.

- **Overall Predictive Performance:** Similar to the synthetic datasets, we report the overall prediction quality using all of the hypotheses in a shared regression, reporting AUC for classification and R^2 for regression.

5.3. Baselines

We compare HYPOTHESAES against recent, state-of-the-art methods (§2). For fair comparison, we run all methods with the same embeddings and LLMs as our method, unless stated otherwise: OpenAI’s text-embedding-3-small for embeddings; GPT-4o for hypothesis proposal; and GPT-4o-mini for hypothesis annotation. We run each method to produce H hypotheses and evaluate all methods identically.

BERTOPIC (Grootendorst, 2022) is a neural topic model which produces topics by clustering text embeddings. To produce H hypotheses from BERTOPIC, texts are clustered into topics, and we fit an L_1 -regularized model to predict Y from topic ID; like our method, we set the penalty such that

⁴We use a Bonferroni-corrected p -value threshold of $2.5e-3$.

there are exactly H nonzero coefficients. We assign each topic a natural language label by prompting an LLM with a sample of documents and top terms (prompt in Figure 7).

NLPARAM (Zhong et al., 2024) trains a neural network with a length- K bottleneck layer to predict Y from text embeddings, and then uses an LLM to propose labels corresponding to each one of the bottleneck neurons’ activations. The hypotheses are iteratively refined; a hypothesis is retained only if its annotations (computed on the entire dataset) improve the loss. We run NLPARAM for the default 10 iterations, with 3 candidate hypotheses per neuron.

HYPOGENIC (Zhou et al., 2024) directly prompts a language model to propose hypotheses given a task-specific prompt and labeled examples. They use a bandit reward function to score hypotheses and propose new ones based on incorrectly-labeled examples. We run HYPOGENIC with default parameters; note that the method uses the same model for proposal and annotation, which defaults to GPT-4o-mini.

Training set size. A key advantage of HYPOTHESAES is that feature generation and selection require only an SAE forward pass, which is cheap even for large datasets; the only LLM-dependent step is feature interpretation. BERTOPIC also scales well, as it only uses LLM inference for topic labeling. In contrast, NLPARAM and HYPOGENIC score and select hypotheses by annotating every training example, so their inference costs scale with training set size. Running either method on 200K Yelp reviews, for instance, would cost $> \$500$ and 78 hours (compared to $\sim \$1$ and 20 min-

METHOD	WIKI, $H = 5$			WIKI, $H = 15$			BILLS, $H = 5$			HEADLINES		YELP		CONGRESS	
	F1	SURF.	AUC	F1	SURF.	AUC	F1	SURF.	AUC	Ct.	AUC	Ct.	R^2	Ct.	AUC
HYPOTHESAES	0.58	0.70	0.84	0.47	0.61	0.84	0.67	0.78	0.89	15	0.69	15	0.76	15	0.70
BERTOPIC	0.48	0.48	0.77	0.40	0.43	0.79	0.65	0.88	0.85	1	0.62	12	0.65	10	0.64
NLPARAM	0.15	0.30	0.69	0.16	0.34	0.77	0.32	0.50	0.74	5	0.66	11	0.58	8	0.65
HYPOGENIC	0.15	0.20	0.60	0.12	0.39	0.64	0.25	0.30	0.70	3	0.65	12	0.78	5	0.68

Table 2. HYPOTHESAES performs best at recovering reference human-annotated topics on WIKI-5, WIKI-15, and BILLS. We also evaluate on three real-world datasets: HEADLINES, YELP, and CONGRESS. “Ct.” refers to the count of significant hypotheses, out of 20 candidates; HYPOTHESAES generates the most significant hypotheses.

utes for HYPOTHESAES). To compare to the baselines as favorably as possible, we allow each method to use up to $50\times$ the cost and runtime of ours. To fit this budget, we use up to 20,000 training samples for NLPARAM and HYPOGENIC, which accommodates the full training sets on WIKI, BILLS, and HEADLINES; we use random subsets for YELP and CONGRESS. Notably, this sample budget is significantly larger than either method uses in their publication experiments (4,748 for NLPARAM; 200 for HYPOGENIC), suggesting our comparisons are generous to the baselines.

6. Results

6.1. Synthetic Datasets

HYPOTHESAES recovers ground-truth hypotheses. Table 2 shows quantitative metrics: HYPOTHESAES beats all baselines on 8 of 9 dataset-metric pairs. As captured by surface similarity, most of HYPOTHESAES’s inferred hypotheses either perfectly match or are related to the references. On WIKI-5 (Table 4), we retrieve two hypotheses exactly; the other three are slightly too specific (HYPOTHESAES infers a topic *is about a New York State Route or highway*, but the reference includes all Northeastern roads, not just New York) or broad (it infers *historical battles* but does not specify *battles since 1800*), but all inferred hypotheses are always related to a reference, i.e., has surface similarity ≥ 0.5 . No baselines meet this standard on WIKI-5: NLPARAM and HYPOGENIC output 3/5 and 2/5 hypotheses that are unrelated to any of the reference topics, respectively. BERTOPIC performs well, but outputs redundant hypotheses about state highways and misses ‘battles’ entirely.

Our hypotheses also closely match ground-truth when used to annotate heldout data. The mean F1 score compares how closely annotations according to an inferred hypothesis (e.g. *is about a Simpsons episode*) match ground-truth labels (whether the WIKI article belongs to the human-annotated Simpsons category). HYPOTHESAES beats all baselines on this metric, especially outperforming NLPARAM (+0.36, on average) and HYPOGENIC (+0.40). Finally, regressing the label using each method’s annotations, HYPOTHESAES achieves higher AUC than all baselines. Overall, HYPOTHESAES usually outperforms BERTOPIC—which is designed

exactly for this topic inference setting—and substantially outperforms two state-of-the-art LLM hypothesis generation methods.

6.2. Real-World Datasets

Table 2 shows quantitative metrics for the real-world datasets. Tables 1, 5, and 6 show all significant hypotheses on HEADLINES, CONGRESS, and YELP respectively.

HYPOTHESAES identifies more distinct hypotheses than other methods. Compared to baselines, HYPOTHESAES produces the most hypotheses which are significant in a multivariate regression. Importantly, this evaluation reflects successful *generalization*: the hypotheses are learned on the training and validation sets, and they remain significantly associated with the target variable when GPT-4o-mini annotates them on the heldout set. Across the three datasets, **45 out of 60 candidate hypotheses are significant**, substantially ahead of the 24, 23, and 20 for NLPARAM, BERTOPIC, and HYPOGENIC respectively.

When evaluating overall predictive performance, HYPOTHESAES beats baselines on 8 of 9 comparisons, demonstrating strongly predictive hypotheses. The exception is that HYPOGENIC achieves 2% higher R^2 on Yelp. However, HYPOGENIC’s most predictive hypotheses, such as *expresses disappointment with the overall dining experience*, essentially restate the rating prediction task. While LLMs can perform this prediction well, these hypotheses do not reveal specific insights for the task. In contrast, our hypotheses—in Table 6 for YELP—tend to be more specific: we find that negative reviews mention food poisoning, rude and aggressive staff, or dishonest business practices.

6.3. Evaluating Hypothesis Novelty

Beyond quantitative comparisons against baselines, we evaluate if HYPOTHESAES discovers new hypotheses even on datasets that have been thoroughly analyzed in prior work.

CONGRESS. A classic study from Gentzkow et al. (2016) identifies bigrams and trigrams in Congressional speeches that mark Republican (R) or Democrat (D) speakers. We compare to their work by computing counts, in each speech,

of their full reported n-gram lists (150 R / 150 D); we also run HYPOTHESAES to produce 100 hypotheses, of which 33 are significant. A regression using only n-gram counts achieves AUC 0.61, while including our 33 hypotheses increases AUC to 0.74. Notably, 28 out of 33 hypotheses remain significantly predictive when controlling for n-grams, showing that our hypotheses add signal to classical methods.

Beyond predictive performance, HYPOTHESAES offers two qualitative advantages. First, our hypotheses are inherently interpretable. For example, while Gentzkow et al. (2016) report that “oil companies” correlates with Democrats, it requires context to interpret: are speakers referring to oil companies positively or negatively? In contrast, HYPOTHESAES’s corresponding hypothesis is *criticizes oil companies or oil industry practices*. Second, our hypotheses capture nuanced patterns that n-grams cannot express, like *criticizes government resource allocation, highlighting disparities between domestic needs and actions abroad*.

HEADLINES. To test our discoveries against more modern methods, we compare to Batista & Ross (2024), an LLM hypothesis generation study designed for HEADLINES and aimed at producing new insights for marketing researchers. Each of their 5 hypotheses is validated via human annotation. Using our heldout evaluation protocol, we compute annotations using their hypotheses, which achieve AUC 0.59 at predicting binary engagement; adding in HYPOTHESAES’s 15 significant hypotheses increases AUC to 0.70. Crucially, **all 15 of our hypotheses remain significant** when controlling for the prior discoveries in a shared regression.

Several other prior works have also studied HEADLINES using classical dictionary-based methods (Robertson et al., 2023; Gligorić et al., 2023; Banerjee & Urminsky, 2024; Aubin Le Quéré & Matias, 2025). HYPOTHESAES qualitatively improves on these findings by producing more specific hypotheses. For example, Robertson et al. (2023) find that negative words increase engagement, and positive words decrease it. HYPOTHESAES supports this finding, and also make it more precise: we find that *situations that evoke discomfort or embarrassment* and mentions of offensive behavior increase clicks, while *initiatives that positively impact a community* decreases them (Table 1). Banerjee & Urminsky (2024) find that words in an “emotional intensity” dictionary increase engagement, which may be difficult to operationalize; our corresponding finding is that headlines referencing *being surprised, shocked, or unprepared* increase engagement. Our findings are more precise, and can therefore provide richer insights to researchers.

These results demonstrate that HYPOTHESAES, despite no dataset-specific design, uncovers novel hypotheses that advance prior work from domain experts—suggesting broad applicability across diverse settings.

6.4. Costs

In Table 3, we report the runtimes, LLM inference token counts, and costs (at current OpenAI API pricing) for all methods on CONGRESS; the trends are similar for all datasets. Runtimes include training the SAEs on one NVIDIA A6000 GPU. This step is fast because our SAEs can afford to be small ($\sim 10^3$ features, compared to 10^7 in Gao et al. (2024)), since our focus is to learn domain-specific features rather than features of all language. HYPOTHESAES is substantially cheaper than NLPARAM and HYPOGENIC: despite the fact that the latter methods are trained on 20K examples on CONGRESS (vs. 114K for HYPOTHESAES), **HYPOTHESAES is $\sim 30\text{-}50\times$ faster and $10\text{-}50\times$ cheaper**. This is because HYPOTHESAES’s annotation requirement scales only with the number of hypotheses H , while these two methods scale with both H and N .

BERTOPIC is cheapest, and remains a good baseline for a resource-limited analysis. However, HYPOTHESAES can also be made cheaper by skipping the non-essential label validation step: instead of generating 3 candidate labels per neuron and choosing the highest-fidelity one, we can simply generate 1 label per neuron and use it directly. This reduces cost to that of BERTOPIC, while still exceeding its performance; we describe this ablation in Appendix B.2.

Method	Time (min)	Tokens (M)	Cost
HYPOTHESAES	18.1	0.2 / 7.2	\$1.29
SAE-NOVAL	7.9	0.1 / 0.0	\$0.15
NLPARAM	904.9	2.8 / 414	\$69.92
HYPOGENIC	575.4	0.0 / 72	\$11.02
BERTOPIC	3.0	0.1 / 0.0	\$0.19

Table 3. Runtimes and costs for hypothesis generation on CONGRESS. We report millions of tokens for GPT-4o (feature labeling) and GPT-4o-mini (annotation) respectively as {# 4o} / {# 4o-mini}. Input and output token counts are summed (>99% are input). SAE-NOVAL is our method without the label validation step (i.e., one label is generated and used per neuron). Costs are in USD and based on OpenAI pricing as of 01/2025.

7. Conclusion

We propose HYPOTHESAES, a scalable method to generate interpretable hypotheses from black-box representations. On three synthetic and three real-world social science tasks, HYPOTHESAES outperforms a topic model and state-of-the-art LLM baselines, at less than $10\times$ the time and cost of the latter. The method applies naturally to many tasks in social science (Ziems et al., 2024), healthcare (Hsu et al., 2023; Robitschek et al., 2025; Pierson et al., 2025), and biology (Vig et al., 2021). More generally, there are many applications where we can accurately *predict* a target variable, but producing new scientific insights remains difficult. Our work helps bridge this gap.

Software and Data

Code is [available on GitHub](#), and can be installed via pip with `pip install hypothesaes`. A notebook to reproduce experimental results in the paper is available on the repository. All data used in the paper are [available on HuggingFace](#). A demo to visualize all SAE neurons and how well they predict the target for HEADLINES, YELP, and CONGRESS is available at <https://hypothesaes.org/>.

Acknowledgements

We thank Rishi Jha and Hal Triedman for helpful discussion about sparse autoencoders, Anna Lyubarskaja for helpful discussion about Proposition 3.1, Ruiqi Zhong for helpful discussion about NLPARAM, Neil Movva and OpenAI for LLM inference credits, Serina Chang for useful comments, Sammi Cheung for writing feedback, and members of the Pierson and Garg groups for comments and discussions.

Impact Statement

Our work advances the field of AI-assisted hypothesis generation, which we hope will drive progress in the social and natural sciences. To the best of our knowledge, there are no particular negative social consequences imposed by our work compared to machine learning research in general.

References

- Aubin Le Qu  r  , M. and Matias, J. N. When curiosity gaps backfire: Effects of headline concreteness on information selection decisions. *Scientific Reports*, 15(1):994, January 2025. ISSN 2045-2322.
- Banerjee, A. and Urminsky, O. The Language That Drives Engagement: A Systematic Large-scale Analysis of Headline Experiments. *Marketing Science*, November 2024. ISSN 0732-2399.
- Batista, R. M. and Ross, J. Words that Work: Using Language to Generate Hypotheses, July 2024.
- Bills, S., Cammarata, N., Mossing, D., Tillman, H., Gao, L., Goh, G., Sutskever, I., Leike, J., Wu, J., and Saunders, W. Language models can explain neurons in language models. <https://openaipublic.blob.core.windows.net/neuron-explainer/paper/index.html>, 2023.
- Blei, D. M., Ng, A. Y., and Jordan, M. I. Latent dirichlet allocation. *J. Mach. Learn. Res.*, 3(null):993–1022, March 2003. ISSN 1532-4435.
- Bricken, T., Templeton, A., Batson, J., Chen, B., Jermyn, A., Conerly, T., Turner, N., Anil, C., Denison, C., Askell, A., et al. Towards monosemanticity: Decomposing language models with dictionary learning. *Transformer Circuits Thread*, 2, 2023.
- Chang, J., Gerrish, S., Wang, C., Boyd-graber, J., and Blei, D. Reading Tea Leaves: How Humans Interpret Topic Models. In *Advances in Neural Information Processing Systems*, volume 22. Curran Associates, Inc., 2009.
- Cunningham, H., Ewart, A., Riggs, L., Huben, R., and Sharkey, L. Sparse autoencoders find highly interpretable features in language models. *arXiv preprint arXiv:2309.08600*, 2023.
- Demszky, D., Garg, N., Voigt, R., Zou, J., Gentzkow, M., Shapiro, J., and Jurafsky, D. Analyzing Polarization in Social Media: Method and Application to Tweets on 21 Mass Shootings, April 2019.
- Elhage, N., Hume, T., Olsson, C., Schiefer, N., Henighan, T., Kravec, S., Hatfield-Dodds, Z., Lasenby, R., Drain, D., Chen, C., Grosse, R., McCandlish, S., Kaplan, J., Amodei, D., Wattenberg, M., and Olah, C. Toy Models of Superposition, September 2022.
- Gao, L., la Tour, T. D., Tillman, H., Goh, G., Troll, R., Radford, A., Sutskever, I., Leike, J., and Wu, J. Scaling and evaluating sparse autoencoders, June 2024.
- Gentzkow, M. and Shapiro, J. M. What Drives Media Slant? Evidence From U.S. Daily Newspapers. *Econometrica*, 78(1):35–71, 2010. ISSN 1468-0262.
- Gentzkow, M., Shapiro, J. M., and Taddy, M. Measuring Group Differences in High-Dimensional Choices: Method and Application to Congressional Speech, July 2016.
- Gligori  , K., Lifchits, G., West, R., and Anderson, A. Linguistic effects on news headline success: Evidence from thousands of online field experiments (Registered Report). *PLOS ONE*, 18(3):e0281682, March 2023. ISSN 1932-6203.
- Grimmer, J. A bayesian hierarchical topic model for political texts: Measuring expressed agendas in senate press releases. *Political analysis*, 18(1):1–35, 2010.
- Grimmer, J., Roberts, M. E., and Stewart, B. M. Machine Learning for Social Science: An Agnostic Approach. *Annual Review of Political Science*, 24(Volume 24, 2021): 395–419, May 2021. ISSN 1094-2939, 1545-1577.
- Grimmer, J., Roberts, M. E., and Stewart, B. M. *Text as Data: A New Framework for Machine Learning and the Social Sciences*. Princeton University Press, January 2022. ISBN 978-0-691-20799-5.

- Grootendorst, M. BERTopic: Neural topic modeling with a class-based TF-IDF procedure, March 2022.
- Grootendorst, M. LLM & Generative AI - BERTopic. https://maartengr.github.io/BERTopic/getting_started/representation/llm.html, 2024.
- Hoyle, A. M., Sarkar, R., Goel, P., and Resnik, P. Are Neural Topic Models Broken? In Goldberg, Y., Kozareva, Z., and Zhang, Y. (eds.), *Findings of the Association for Computational Linguistics: EMNLP 2022*, pp. 5321–5344. Association for Computational Linguistics, December 2022.
- Hsu, C.-C., Obermeyer, Z., and Tan, C. Clinical Notes Reveal Physician Fatigue, December 2023.
- Kim, B., Wattenberg, M., Gilmer, J., Cai, C., Wexler, J., Viegas, F., and Sayres, R. Interpretability Beyond Feature Attribution: Quantitative Testing with Concept Activation Vectors (TCAV). In *Proceedings of the 35th International Conference on Machine Learning*, pp. 2668–2677. PMLR, July 2018.
- Ludan, J. M., Lyu, Q., Yang, Y., Dugan, L., Yatskar, M., and Callison-Burch, C. Interpretable-by-Design Text Understanding with Iteratively Generated Concept Bottleneck, April 2024.
- Ludwig, J. and Mullainathan, S. Machine Learning as a Tool for Hypothesis Generation, March 2023.
- Makhzani, A. and Frey, B. K-Sparse Autoencoders, March 2014.
- Matias, J. N., Munger, K., Le Quere, M. A., and Ebersole, C. The Upworthy Research Archive, a time series of 32,487 experiments in U.S. media. *Scientific Data*, 8(1): 195, August 2021. ISSN 2052-4463.
- Merity, S., Xiong, C., Bradbury, J., and Socher, R. Pointer Sentinel Mixture Models, September 2016.
- Monroe, B. L., Colaresi, M. P., and Quinn, K. M. Fightin’ Words: Lexical Feature Selection and Evaluation for Identifying the Content of Political Conflict. *Political Analysis*, 16(4):372–403, August 2009. ISSN 1047-1987, 1476-4989.
- Nelson, L. K. Computational grounded theory: A methodological framework. *Sociological Methods & Research*, 49(1):3–42, 2020.
- Olah, C., Mordvintsev, A., and Schubert, L. Feature Visualization. *Distill*, 2(11):10.23915/distill.00007, November 2017. ISSN 2476-0757.
- O’Neill, C., Ye, C., Iyer, K., and Wu, J. F. Disentangling Dense Embeddings with Sparse Autoencoders, August 2024.
- Pham, C. M., Hoyle, A., Sun, S., Resnik, P., and Iyyer, M. TopicGPT: A Prompt-based Topic Modeling Framework, April 2024.
- Pierson, E., Shanmugam, D., Movva, R., Kleinberg, J., Agrawal, M., Dredze, M., Ferryman, K., Gichoya, J. W., Jurafsky, D., Koh, P. W., Levy, K., Mullainathan, S., Obermeyer, Z., Suresh, H., and Vafa, K. Using Large Language Models to Promote Health Equity. *NEJM AI*, 2(2):AIp2400889, January 2025.
- Ranard, B. L., Werner, R. M., Antanavicius, T., Schwartz, H. A., Smith, R. J., Meisel, Z. F., Asch, D. A., Ungar, L. H., and Merchant, R. M. Yelp reviews of hospital care can supplement and inform traditional surveys of the patient experience of care. *Health Affairs*, 35(4):697–705, 2016.
- Rathje, S., Van Bavel, J. J., and Van Der Linden, S. Out-group animosity drives engagement on social media. *Proceedings of the National Academy of Sciences*, 118(26): e2024292118, 2021.
- Robertson, C. E., Pröllochs, N., Schwarzenegger, K., Pärnamets, P., Van Bavel, J. J., and Feuerriegel, S. Negativity drives online news consumption. *Nature Human Behaviour*, 7(5):812–822, May 2023. ISSN 2397-3374.
- Robitschek, E., Bastani, A., Horwath, K., Sordean, S., Pletcher, M. J., Lai, J. C., Galletta, S., Ash, E., Ge, J., and Chen, I. Y. A large language model-based approach to quantifying the effects of social determinants in liver transplant decisions, January 2025.
- Sun, C.-E., Oikarinen, T., Ustun, B., and Weng, T.-W. Concept Bottleneck Large Language Models, December 2024.
- Sun, M., Oliwa, T., Peek, M. E., and Tung, E. L. Negative Patient Descriptors: Documenting Racial Bias In The Electronic Health Record. *Health Affairs*, 41(2):203–211, February 2022. ISSN 0278-2715.
- Templeton, A., Conerly, T., Marcus, J., Lindsey, J., Bricken, T., Chen, B., Pearce, A., Citro, C., Ameisen, E., Jones, A., et al. Scaling monosemanticity: Extracting interpretable features from claude 3 sonnet. transformer circuits thread, 2024.
- Tibshirani, R. Regression Shrinkage and Selection Via the Lasso. *Journal of the Royal Statistical Society: Series B (Methodological)*, 58(1):267–288, January 1996. ISSN 0035-9246.

- Ting, P.-J. L., Chen, S.-L., Chen, H., and Fang, W.-C. Using big data and text analytics to understand how customer experiences posted on yelp. com impact the hospitality industry. *Contemporary Management Research*, 13(2), 2017.
- Vig, J., Madani, A., Varshney, L. R., Xiong, C., Socher, R., and Rajani, N. F. BERTology Meets Biology: Interpreting Attention in Protein Language Models, March 2021.
- Yelp. Yelp Open Dataset. <https://business.yelp.com/data/resources/open-dataset/>, 2024.
- Zhong, R., Wang, H., Klein, D., and Steinhardt, J. Explaining Datasets in Words: Statistical Models with Natural Language Parameters, September 2024.
- Zhou, Y., Liu, H., Srivastava, T., Mei, H., and Tan, C. Hypothesis Generation with Large Language Models, August 2024.
- Ziems, C., Held, W., Shaikh, O., Chen, J., Zhang, Z., and Yang, D. Can Large Language Models Transform Computational Social Science?, February 2024.

Sparse Autoencoders for Hypothesis Generation

REFERENCE	INFERRED	F1	SURFACE
Other music articles > Performers, groups, composers, and other music-related people	<i>is about a rock band or a musical act</i>	0.39	0.5
Earth science > Tropical cyclones: Atlantic	<i>is about hurricanes or tropical storms</i>	0.52	1.0
Battles, exercises, and conflicts > Modern history (1800 to present)	<i>is about a historical battle</i>	0.36	0.5
Television > The Simpsons episodes	<i>is about episodes of The Simpsons television show</i>	0.92	1.0
Transport > Road infrastructure: Northeastern United States	<i>is about a New York State Route</i>	0.71	0.5

Table 4. Evaluating hypotheses inferred by HYPOTHESAES against the reference topics on WIKI.

Source	Overall AUC	# Sig	Significant Hypotheses: Red: ↑ Republican ; Blue: ↑ Democrat	Sep.	AUC
HYPOTHESAES	0.702	15	<i>contains the phrase 'I ask unanimous consent'</i>	0.30	0.53
			<i>mentions hearings conducted by Senate Committees or Subcommittees</i>	0.18	0.52
			<i>mentions scheduling or details about Senate votes or amendments</i>	0.17	0.55
			<i>discusses illegal immigration and its associated implications</i>	0.15	0.51
			<i>mentions victories or successes in military or political contexts</i>	0.15	0.51
			<i>mentions freedom or liberty</i>	0.06	0.50
			<i>discusses economic growth or growth rates</i>	0.01	0.50
			<i>discusses government spending or budgetary issues</i>	-0.09	0.54
			<i>mentions the need for reform or proposes reforms</i>	-0.17	0.58
			<i>mentions key figures related to the civil rights movement</i>	-0.22	0.51
			<i>criticizes government policies or actions</i>	-0.29	0.63
			<i>mentions the U.S. national debt or raising the debt limit</i>	-0.40	0.51
			<i>criticizes tax breaks or advantages for the wealthy</i>	-0.44	0.52
			<i>criticizes Republican leadership or policies</i>	-0.45	0.60
			<i>criticizes the administration's handling of the Iraq war</i>	-0.45	0.52
BERTOPIC	0.636	10	<i>Senate procedural requests and adjournment schedules</i>	0.39	0.53
			<i>Requests for unanimous consent to authorize committee meetings</i>	0.32	0.53
			<i>procedural discussions and scheduling of votes</i>	0.20	0.56
			<i>discusses immigration and border security</i>	0.13	0.51
			<i>mentions individuals with professional titles or roles</i>	0.08	0.53
			<i>debates on earmarks and federal spending processes</i>	0.05	0.50
			<i>discusses the Darfur conflict and international intervention</i>	-0.30	0.50
			<i>Congressional oversight of defense contracting in Iraq</i>	-0.33	0.50
			<i>discusses Asian Pacific American communities</i>	-0.36	0.50
NLPARAM	0.650	8	<i>Critiques Republican majority leadership</i>	-0.46	0.55
			<i>expresses strong patriotism</i>	0.09	0.51
			<i>contains specific legislative references</i>	0.08	0.53
			<i>mentions specific individuals or organizations</i>	-0.03	0.52
			<i>includes numerical data</i>	-0.08	0.53
			<i>addresses a critical national issue</i>	-0.14	0.56
			<i>expresses strong support or opposition</i>	-0.15	0.58
			<i>advocates for specific groups or communities</i>	-0.17	0.55
HYPOGENIC	0.675	5	<i>employs emotional or dramatic language</i>	-0.23	0.58
			<i>focuses on national security and immigration enforcement</i>	0.18	0.51
			<i>discusses social issues such as healthcare and civil rights</i>	-0.22	0.58
			<i>mentions of government inefficiencies or calls for reform</i>	-0.27	0.62
			<i>advocates for civil liberties and critiques government overreach</i>	-0.29	0.57
			<i>emphasizes the need for government accountability</i>	-0.37	0.62

Table 5. CONGRESS hypotheses, for each method, with $p < 0.05$ after Bonferroni correction in a multivariate regression. The ‘Sep.’ column shows the separation score: $E[Y|\hat{Z} = 1] - E[Y|\hat{Z} = 0]$, where positive scores predict Republican speech and negative scores predict Democratic speech. The ‘AUC’ column provides univariate AUCs for individual hypotheses. ‘Overall AUC’ uses all hypotheses, including insignificant ones.

Source	Overall R^2	# Sig	Significant Hypotheses: Green: $\uparrow$ rating ; Red: $\downarrow$ rating	Sep.	R^2
HYPOTHESAES	0.755	15	<i>expresses extreme enthusiasm or excitement with phrases like 'holy fuck', 'can't wait to go back', or extensive use of exclamation points</i>	1.45	0.23
			<i>uses exclamation marks enthusiastically to convey positive sentiment</i>	1.41	0.23
			<i>mentions anticipation or plans to return to the restaurant or order again</i>	1.20	0.17
			<i>explicitly states that a food item or restaurant is the best, often using superlative language</i>	1.10	0.07
			<i>mentions mediocrity or averageness of food or experience</i>	-1.84	0.31
			<i>mentions issues related to food safety, hygiene, or improper food handling</i>	-1.98	0.15
			<i>mentions issues with restaurant staff communication or behavior</i>	-2.23	0.44
			<i>mentions dissatisfaction or criticism of service, staff, or management</i>	-2.25	0.61
			<i>mentions food poisoning or symptoms of foodborne illness</i>	-2.32	0.02
			<i>mentions rude or aggressive behavior from restaurant staff</i>	-2.37	0.30
			<i>mentions dissatisfaction with restaurant policies, rules, or restrictions</i>	-2.37	0.53
			<i>criticizes the behavior or attitude of staff</i>	-2.40	0.50
			<i>complains about rude or unprofessional behavior from staff or management</i>	-2.42	0.46
			<i>mentions dishonest or unethical business practices by the establishment</i>	-2.48	0.17
			<i>describes a specific instance of poor customer service</i>	-2.50	0.50
BERTOPIC	0.647	12	<i>praises the combination of excellent food quality and friendly, attentive service</i>	1.61	0.33
			<i>describes detailed experiences with food quality, preparation, and flavor nuances</i>	0.46	0.03
			<i>criticizes the freshness, quality, or preparation of sushi and fish dishes</i>	-1.77	0.07
			<i>describes issues with restaurant closures or inaccurate operating hours</i>	-1.77	0.10
			<i>complains about uncleanliness of restaurant spaces</i>	-1.91	0.05
			<i>criticizes the quality and authenticity of Chinese dishes</i>	-1.94	0.06
			<i>criticizes Italian dishes, particularly pasta and sauces</i>	-1.94	0.04
			<i>complains about issues with pizza orders</i>	-2.08	0.16
			<i>criticizes breakfast food quality, portion sizes, and service experience</i>	-2.15	0.37
			<i>complains about issues with delivery orders</i>	-2.20	0.25
NLPARAM	0.583	11	<i>uses enthusiastic or excessively positive language</i>	1.90	0.47
			<i>describes exceptional food quality</i>	1.65	0.35
			<i>describes a positive recurring experience</i>	1.45	0.26
			<i>recommends the restaurant to others</i>	1.29	0.18
			<i>uses exclamation marks to express enthusiasm</i>	1.22	0.17
			<i>mentions attentive or friendly staff</i>	1.01	0.12
			<i>expresses enthusiasm for the ambiance or atmosphere</i>	0.92	0.07
			<i>mentions thoughtful accommodation of customer preferences</i>	0.79	0.01
			<i>describes generosity in portion sizes</i>	0.69	0.02
			<i>comments on portion size</i>	0.02	0.00
HYPOGENIC	0.775	12	<i>mentions long wait times to receive service</i>	-1.70	0.15
			<i>highlights a pleasant atmosphere or enjoyable dining experience</i>	2.40	0.66
			<i>expresses a willingness to return or recommend the restaurant to others</i>	2.13	0.55
			<i>expresses a desire to return or recommends the place to others</i>	1.89	0.45
			<i>describes food as delicious or highlights specific dishes with positive adjectives</i>	1.83	0.42
			<i>notes issues with food quality, such as being cold or poorly cooked</i>	-1.87	0.29
			<i>includes complaints about long wait times for food or service</i>	-1.94	0.30
			<i>expresses disappointment with the quality of food compared to expectations</i>	-2.17	0.49
			<i>mentions poor service or unresponsive staff during the dining experience</i>	-2.36	0.51
			<i>complains about poor customer service or rude staff</i>	-2.41	0.50
			<i>complains about poor service or unhelpful staff</i>	-2.42	0.60
			<i>expresses disappointment with the overall dining experience</i>	-2.42	0.67
			<i>expresses disappointment with the overall experience or ambiance</i>	-2.43	0.67

Table 6. YELP hypotheses, for each method, with $p < 0.05$ after Bonferroni correction in a multivariate regression. ‘Sep.’ is the separation score: $E[Y|\hat{Z} = 1] - E[Y|\hat{Z} = 0]$, where positive scores predict higher restaurant ratings. The ‘ R^2 ’ column provides univariate R^2 values of individual hypotheses. ‘Overall R^2 ’ uses all 20 hypotheses, including insignificant ones. Hypotheses for all methods are abbreviated.

A. Data Preprocessing

WIKI. We download the dataset as processed by [Pham et al. \(2024\)](#); the data are derived from the WikiText corpus assembled by [Merity et al. \(2016\)](#). The articles are categorized by Wikipedia editors; each article has a supercategory, category, subcategory (e.g., Media and Drama > Television > The Simpsons Episodes). We focus on recovering article subcategories, since these are the most specific and therefore most difficult to fully recover. After filtering out duplicates and infrequent (<100 articles) subtopics, the dataset contains 11,979 items spanning 69 subcategories. For the WIKI-5 dataset, we set the ground-truth topics to the 5 most common subcategories. That is, we label all 1,771 articles (14.8%) in these subcategories as positives, while all other article subcategories are negatives. The WIKI-15 dataset is constructed similarly using the 15 most common subcategories (4,190 positives, 35.0%).

BILLS. We also download BILLS from [Pham et al. \(2024\)](#); the dataset was originally assembled by [Hoyle et al. \(2022\)](#) using GovTrack⁵. The bills are from the 110th-114th U.S. Congresses (Jan 2007 - Jan 2017). Each bill has a topic and subtopic. After filtering out very short bill texts (which may not have been properly transcribed), duplicate bill texts, bills with no subtopic, and bills with infrequent (<100) subtopics, there are 20,834 items spanning 70 subtopics. We label the 5,038 bills (24.2%) with the 5 most common subtopics as positives, and all others as negatives.

HEADLINES. We use the HEADLINES dataset as collected and released by the Upworthy Research Archive ([Matias et al., 2021](#)). All of the data is derived from web traffic to Upworthy.com, a digital media platform known for its attention-grabbing articles. The dataset consists of thousands of A/B tests, where users were randomly shown one of several possible headline variations for the same underlying article. Each headline is linked to a corresponding number of impressions and clicks, allowing researchers to estimate the causal effect of headline text on *click-rate* (clicks/impressions, also referred to as clickthrough-rate or engagement).

To preprocess the data, we start by downloading all publicly available data from the Upworthy archive⁶, which includes clicks and impressions for 122,565 distinct headlines (we remove 28K rows from a problematic data range as recently reported by the dataset authors⁷). We then group headlines which have the same `test_id` (signifying that those headlines were randomized against one another). There are 14,128 such groups with at least two headlines. Following [Zhou et al. \(2024\)](#), we take the headlines with max and min click-rates in each group, and design a binary classification task to predict which headline in the pair received higher engagement. That is, we construct tuples (headline A, headline B, label), where headline A and headline B are randomly shuffled, and the label is 1 if headline A received higher click-rate than headline B.

Following [Batista & Ross \(2024\)](#), we take specific care to prevent train-holdout leakage on this dataset. Specifically, the same headline texts may be used in multiple randomized experiments (i.e., the same headline may occur in different `test_id`'s). We ensure that no identical headlines will appear in both the train split and the heldout split; we encourage future researchers to follow this practice, since splitting on `test_id` alone does not fully prevent headline leakage.

To handle the pairwise design, we train the SAE on the embeddings of all unique headlines in the training set, and then compute the activation matrices Z_{SAE}^A and Z_{SAE}^B for all of the headline A's and headline B's respectively. We select predictive neurons by identifying columns in the matrix difference $\Delta Z_{SAE} := Z_{SAE}^A - Z_{SAE}^B$ which are predictive of y .

YELP. We download the Yelp Open Dataset⁸, and we filter to the 4.72M reviews where the business contains the 'Restaurant' tag. From this, we randomly sample training, validation, and heldout sets (200K, 10K, and 10K respectively).

CONGRESS. We download all congressional speech transcripts from the 109th U.S. Congress ([Gentzkow & Shapiro, 2010](#))⁹, which met from January 2005 to January 2007. We filter speeches which are very short (less than 250 characters) and include only speeches from Republican or Democrat speakers; the dataset is roughly balanced (51.7% Republican). If a speech is longer than ten sentences, we split it into ten sentence chunks, and exclude any chunks beyond the first five (to avoid over-representation of specific speeches). For the heldout set, we restrict to one chunk per speech to ensure that our hypotheses generalize to a diverse corpus. Ultimately, we use 114K speech chunks for training, 16K for validation, and 12K for heldout eval. If a speech has multiple excerpts, we ensure that all of them are included in the same split to avoid leakage.

⁵<https://www.govtrack.us/congress/bills/>

⁶<https://upworthy.natematias.com/index>

⁷<https://upworthy.natematias.com/2024-06-upworthy-archive-update.html>

⁸<https://business.yelp.com/data/resources/open-dataset/>

⁹https://data.stanford.edu/congress_text

B. Hyperparameters

B.1. Training sparse autoencoders

We train k -sparse autoencoders, where each forward pass masks all neurons besides the top- k activating ones (Makhzani & Frey, 2014; Gao et al., 2024). The most influential hyperparameters, which we tune for each dataset, are M and k . We choose (M, k) by maximizing validation performance. Specifically, for each combination of (M, k) , we evaluate validation AUC (or R^2 on YELP) after fitting an L_1 -regularized predictor with exactly $H \ll M$ nonzero coefficients. We constrain M and k to powers of two to tractably limit the search space.

These parameters, broadly, control the level of granularity of the concepts learned by each SAE neuron: M is the total number of concepts which can be learned across the entire dataset, and k is the number of concepts which can be used to represent each instance. For intuition, assume that with $(M, k) = (16, 4)$ on YELP, the SAE learns a single neuron that fires when reviews mention *price*; with $(M, k) = (32, 4)$, the SAE may learn separate neurons for reviews which mention *high prices* vs. *low prices* (described as “feature splitting” in Bricken et al. (2023)); with $(M, k) = (32, 8)$, the SAE may learn to represent more niche features altogether, such as the *number of exclamation marks*. In general, increasing either parameter yields more granular features, at the risk of producing some neurons that are uninterpretable or redundant.

In some cases we would like features at multiple levels of granularity: for example, we may want both the general feature “mentions prices” and the more specific feature “mentions that the cocktails were cheap during happy hour”, but training a single large SAE (e.g. $(1024, 32)$) may only produce features as specific as the latter. A simple solution is to train multiple SAEs of different sizes, concatenate their activation matrices, and select features from any of them. We find that this approach is helpful (improves validation AUC) on the HEADLINES dataset. Ultimately, we use the following values:

- WIKI-5: $(32, 4)$.
- WIKI-15 & BILLS: $(64, 4)$.
- HEADLINES: $(256, 8); (32, 4)$.
- YELP: $(1024, 32)$.
- CONGRESS: $(4096, 32)$.

Following prior work, we add several optimizations to improve SAE training. Following best practices based on results from Bricken et al. (2023), we (i) tie the weights of the pre-encoder bias and the decoder bias, $b_{\text{pre}} = b_{\text{dec}}$; (ii) initialize b_{pre} to the geometric median of the training set; (iii) initialize $W_{\text{dec}} = W_{\text{enc}}^T$; and (iv) force all decoder rows to have unit norm.

To minimize the number of “dead” latent neurons in the SAE—neurons which are never selected by the TopK filter—we use an auxiliary loss (Gao et al., 2024; O’Neill et al., 2024). The auxiliary loss term provides nonzero gradients for up to k_{aux} neurons which were not active in the last several forward passes. We fit these neurons to the residual reconstruction error left by the TopK neurons. That is, let z_i^{AuxK} consist of the activations of the top k_{aux} dead neurons in z_i ; then

$$\mathcal{L}_{\text{aux}} = \|W_{\text{dec}} z_i^{\text{AuxK}} - (e_i - \hat{e}_i)\|_2^2,$$

and the full loss, then, is $\mathcal{L} = \mathcal{L}_{\text{TopK}} + w_{\text{aux}} \mathcal{L}_{\text{AuxK}}$.

Hyperparameters for the auxiliary loss and other training hyperparameters are given below (identical for all experiments):

- k_{aux} : $2 \cdot k$; that is, up to $2 \cdot k$ dead neurons are used to fit the TopK reconstruction error.
- **Dead neuron threshold steps**: 256. This is the number of consecutive steps for which a neuron was not selected by TopK to be considered “dead.”
- **Auxilliary loss coefficient** w_{aux} : $1/32$. This parameter gives the weight of the auxiliary reconstruction’s loss.
- **Batch size**: 512; **learning rate**: $5e-4$; **gradient clipping threshold**: 1.0.
- **Epochs**: Up to 200, with early stopping after 5 epochs of validation loss not decreasing.

B.2. Interpreting neurons with LLMs

There are several parameters for neuron interpretation which we fix across experiments:

- **Interpretation model:** GPT-4o (version 2024-11-20).
- **Temperature:** 0.7.
- **Number of highly-activating examples:** 10.
- **Number of weakly-activating examples:** 10.
- **Maximum word count per example:** 256 (examples longer than this are truncated).
- **Number of candidate interpretations:** 3 (we choose the highest-fidelity interpretation out of 3 candidates).
- **Number of samples to evaluate fidelity:** 200.

In Appendix C, we describe experiments which justify these choices: they either improve fidelity or have no effect compared to alternate choices. We also provide the full interpretation prompt in Figure 4.

There is one hyperparameter which we adjust per dataset: the **bin percentiles** from which we sample highly-activating and weakly-activating examples for the neuron interpretation prompt. By default, highly-activating texts are sampled from the top decile of positive activations [90, 100], and weakly-activating texts are sampled from the bottom decile of positive activations [0, 10]. This works well in general, especially for larger datasets (YELP, CONGRESS). For smaller datasets like WIKI, the [90, 100] bin produces interpretations which are too specific (e.g., “articles about Napoleon”, even if the neuron fires on articles about European leaders more generally). Note that this requires a judgment call: what is “too specific” depends on practitioner needs. For example, on WIKI, we choose bins such that the resulting interpretations are roughly as granular as the topics in the dataset¹⁰. We ultimately choose the following percentile bins for highly-activating texts:

- WIKI-5, WIKI-15, HEADLINES: [80, 100].
- BILLS: [50, 100].
- YELP, CONGRESS: [90, 100].

Choosing the **number of candidate interpretations** depends on budget. In Appendix C we find that selecting the interpretation with highest fidelity from a pool of 3 candidates increases the resulting fidelity modestly, yet statistically significantly; we therefore use **3 candidate interpretations** for all experiments. However, we also ablate this step—that is, generating a single candidate interpretation and using it directly, without computing fidelity. On CONGRESS (we did not test on the other datasets), the resulting hypotheses are of similar quality: 14/20 are significant, with AUC 0.70, compared to 15/20 significant and AUC 0.70 when using 3 candidate interpretations¹¹. This variant of our method, SAE-NOVAL, costs as little as BERTOPIC (Table 3).

C. Supporting Experiments

Our final neuron interpretation procedure for our main experiments is described in §4, with hyperparameters provided in Appendix B.2. Here, we describe several experiments to justify our procedure. The section is organized as follows. In C.1, we show that our metric of interpretation fidelity—the F1 score comparing concept annotations against neuron activations—is a useful selection metric to improve downstream hypothesis quality. In C.2, we test several different interpretation strategies and hyperparameters and measure their impacts on fidelity. In C.3, we analyze how much prediction signal is lost at each stage of our method, shedding light on how future work can improve hypothesis quality.

¹⁰Choosing this parameter requires similar decision-making to choosing the minimum cluster size in BERTOPIC.

¹¹To illustrate this point, we tested solely on CONGRESS; we expect similarly modest differences on other datasets.

C.1. Higher-fidelity interpretations improve hypotheses.

To test whether fidelity is a useful metric to select from different neuron interpretations, we conduct the following experiment on the Yelp dataset:

1. Train an SAE on Yelp (Step 1), and identify the top-20 predictive neurons via Lasso (Step 2).
2. Generate 10 candidate interpretations (at temperature 0.7, and with different random samples in the interpretation prompt) for each neuron using the hyperparameters in B.2 (Step 3).
3. Compute the *fidelity* for each candidate interpretation: that is, binarizing 100 examples in the top decile of the neuron’s positive activations as positives, and 100 examples in the bottom decile of the neuron’s positive activations as negatives, compute the F1 score between the GPT-4o-mini annotations according to the interpretation against the binarized neuron activations.
4. Store two sets of neuron interpretations: one set includes the interpretations with the best F1 score for each neuron (high-fidelity), and the other set includes the interpretations with the worst F1 score for each neuron (low-fidelity).
5. Compute annotations on the full heldout set (10K examples) for the high-fidelity hypotheses and the low-fidelity hypotheses, and compare how predictive they are as measured by R^2 .

This procedure yields three interesting results. First, there is enough variation in the 10 candidate interpretations such that the best and worst F1 scores are substantially different: the median best F1 is 0.79, while the median worst F1 is 0.53. Second, we find that the high-fidelity set is indeed more predictive: in a multivariate OLS, the high-fidelity set achieves an R^2 of 0.76, compared to 0.61 for the low-fidelity set.

Third, most importantly, using a neuron’s high-fidelity interpretation results in a more predictive hypothesis, on average. Figure 2 shows that, for 17 out of 20 neurons, moving from the low-fidelity interpretation to the high-fidelity interpretation increases the R^2 of the resulting hypothesis (p -value, paired t -test: 0.002). For example, for one neuron, the high-fidelity interpretation is “*mentions disputes or issues related to billing or refunds*” (F1: 0.79), while the low-fidelity interpretation is more generically “*mentions an attempt by the restaurant or management to address or respond to a problem or complaint.*” The former hypothesis achieves heldout $R^2 = 0.10$ with restaurant rating, while the latter is not significantly correlated with rating ($R^2 = 0.003$). These experiments show that **fidelity—as measured via F1 score—is a useful heuristic to choose between multiple candidate neuron interpretations**, when our goal is to produce predictive hypotheses.

Finally, following prior work (Bills et al., 2023), we also tested another metric of fidelity: the R^2 between the concept annotations and the neuron activations across 100 top-activating samples and 100 random samples. Choosing the top candidate interpretation based on this alternative fidelity metric did not produce hypotheses that were significantly more or less predictive than using our F1 fidelity metric.

C.2. Effects of the neuron interpretation procedure on fidelity.

Having established that improving interpretation fidelity improves downstream predictiveness, we search across different hyperparameters for neuron interpretation to maximize fidelity. Our final procedure is to generate 3 candidate interpretations for each neuron by prompting GPT-4o with 10 examples in the top decile of positive activations vs. 10 examples in the bottom decile of positive activations. We use the highest fidelity interpretation out of these 3. There are two sources of non-determinism: each candidate interpretation is generated with a different random seed to sample examples, and we use an LLM temperature of 0.7.

To compare this strategy against other strategies—or to compare strategy A and strategy B more generally—we compute interpretations, and their F1 scores, for the top 100 neurons on YELP under both strategies separately. We run a t -test comparing the pairs of resulting F1 scores to test whether one strategy produces consistently higher F1 scores than the other.

We observe the following:

- **Number of candidate interpretations:** Relative to baseline (generating 3 candidate interpretations and taking the one with highest F1 score), generating only a single candidate interpretation is significantly worse (-0.04). Generating 5 candidate interpretations improves significantly on 3 ($+0.01$ F1, paired t -test $p = 0.003$), but the improvement is modest and results in a slower, more expensive method. Therefore, we use only 3 candidate interpretations.

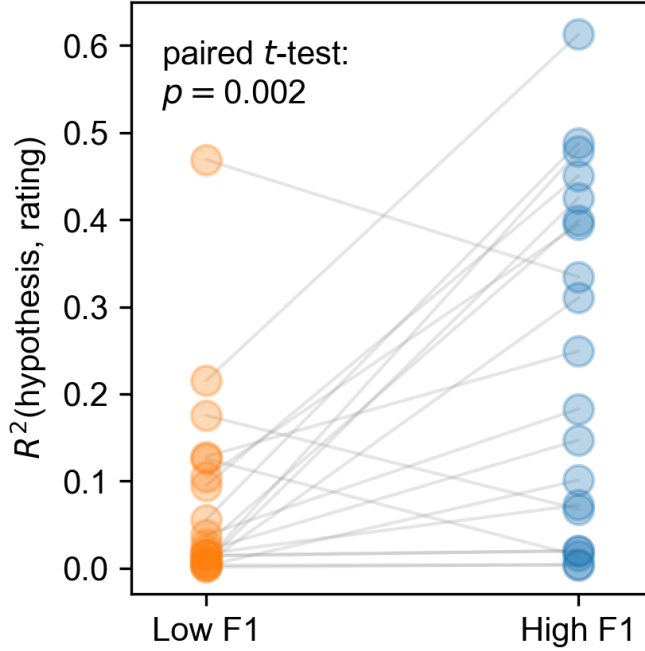


Figure 2. For 17 of the top 20 neurons on YELP, a higher-fidelity neuron interpretation results in a more predictive hypothesis. Fidelity is measured by the F1 score between the concept annotations and the neuron activations on a sample of 100 highly-activating and 100 weakly-activating examples. We generate 10 candidate interpretations for each neuron, and compare the interpretations with best- and worst-F1. Predictiveness is measured by computing the R^2 between the concept annotations and Y (restaurant rating) for 10K heldout examples.

- **Top-random vs. high-weak sampling strategy:** Prior work (e.g. Templeton et al. (2024); Bills et al. (2023)) interprets SAE neurons by prompting with the top-activating texts and random texts. Relative to our baseline (prompting with highly-activating texts, but not necessarily top-activating; and weakly-activating texts instead of random), this alternate approach has no significant effect on F1 score, though the interpretations tend to be more specific: that is, they have higher precision and lower recall. This is because a neuron’s absolute top-activating examples generally share a more specific characteristic than examples sampled randomly from the top decile. Though neither strategy wins clearly in terms of F1, we choose high-weak sampling for our experiments because, empirically, top-random produces hypotheses which can be *too* specific (thereby lowering predictiveness); this is discussed further in B.2. **Our code defaults to top-random** for simplicity, but practitioners should treat sampling strategy as a tunable hyperparameter that depends on desired hypothesis granularity.
- **Continuous vs. binary prompt structure:** Prior work (e.g. Zhong et al. (2024)) prompts the interpreter LLM with a sorted list of texts and their continuous activations. Relative to our baseline (showing separate lists of positive samples and negative samples), this significantly lowers F1 score (-0.06), so we use the binary prompt structure.
- **Chain-of-thought:** Prior work (e.g. O’Neill et al. (2024)) uses a chain-of-thought reasoning prompt to interpret neurons, where the LLM is encouraged to iteratively discover a concept which all of the positive samples contain and none of the negative samples contain. Relative to baseline, we found that CoT has an insignificant but weakly negative effect (-0.03 mean F1), and also sometimes fails to output any interpretation. However, this is still a promising line for future work: it’s possible that while GPT-4o does not reason well, better models or reasoning-specific models (OpenAI o1, DeepSeek R1, etc.) will produce higher fidelity labels with chain-of-thought.
- **Sample count:** Relative to baseline (prompting with 20 total samples, 10 highly-activating and 10 weakly-activating), prompting with 10 total samples lowers average F1 score by 0.03. Prompting with 50 samples has no effect, so we use 20 to favor shorter prompts.
- **Temperature:** Relative to baseline (temperature 0.7), using a lower (0.0) or higher (1.0) temperature has an insignificant but weakly negative effect (-0.01 in mean F1), so we use 0.7.

C.3. Localizing losses in predictive performance to different steps of the method.

HYPOTHESAEs starts from embeddings and ultimately produces hypotheses, after three steps. It is worth asking—how does predictive performance change at each step of our method? Answering this question points to possible improvements: if the SAE’s full hidden representation is much less predictive of restaurant ratings than the input text embeddings, perhaps we need a larger (more expressive) SAE, or more data to train the SAE’s representations; meanwhile, if the largest performance loss comes from translating neurons into natural language hypotheses, perhaps we need a stronger LLM for interpretation.

Figure 3 displays this experiment. At each stage, we fit a logistic or linear regression on the training set to predict the labels on HEADLINES, YELP, and CONGRESS. The four stages are: (i) text embeddings (length-1536 vectors of floats); (ii) the full SAE activation matrix, Z_{SAE} (length-288, 1024, and 4096 vectors of floats for the three tasks, respectively); (iii) the activations of the top-20 selected SAE neurons (length-20 vectors of floats); and (iv) the annotations from the interpreted hypotheses (length-20 vectors of binary 0/1 values).

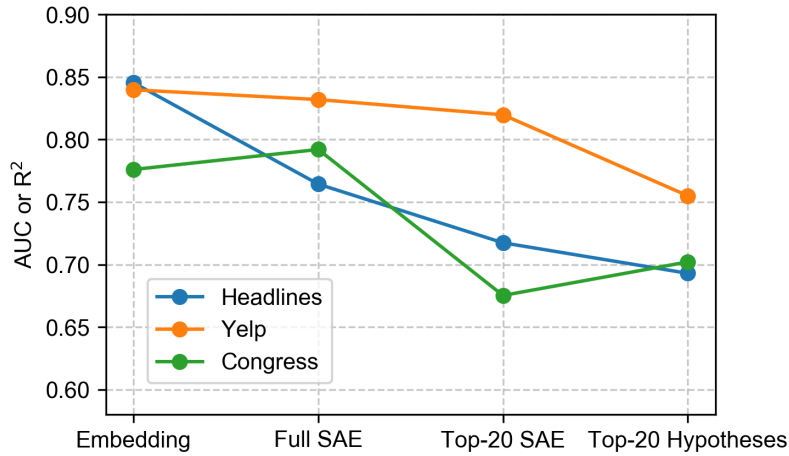


Figure 3. Performance losses at each step of the HYPOTHESAEs pipeline. The leftmost point shows the performance of training a supervised single-layer model to predict the outcome using text embeddings; the second point shows the supervised performance using the full SAE representation; the third point shows the supervised performance using the top-20 SAE neurons; the rightmost point shows the supervised performance using the top-20 hypothesis annotations after neuron interpretation.

The trends depend on the dataset. On HEADLINES, we find that the performance drops most at the first stage (-0.08 AUC): the full SAE is significantly less expressive than the original embeddings. This could suggest using a larger SAE, but empirically we find that increasing the SAE size does not improve AUC, perhaps because there is not enough data (~14K headlines) to make use of more parameters. On YELP, we find that performance is mostly retained until neuron interpretation, at which point R^2 drops by 0.06. This suggests some combination of: our neuron interpretations are imperfect, GPT-4o-mini is not annotating the hypotheses faithfully, or binary annotations are not expressive enough to capture the neuron activations (which are continuous). Finally, on CONGRESS, the only drop comes from using the top-20 neurons to predict the label instead of the full SAE representation (4096 neurons). This implies that 20 features aren’t expressive enough; or, more optimistically, that there are more than 20 distinct hypotheses to discover on this dataset (see §6.3). Remarkably, there are no losses in the interpretation step on the CONGRESS dataset, suggesting that on this dataset our interpretations are very high-fidelity. Ultimately, this analysis can be a useful diagnostic tool to direct one’s efforts when trying to improve hypothesis quality: losses from embedding $\Rightarrow$ full SAE suggest that the dataset and/or the SAE are too small; losses from full SAE $\Rightarrow$ top- H SAE suggest that H is too small; losses from top- H SAE $\Rightarrow$ top- H hypotheses suggest that the interpretation step can be improved. Meanwhile, if the top- H hypotheses achieve similar AUC to the underlying embeddings, this implies that there may be no more “juice to squeeze” in the dataset.

D. Hyperparameters for Baseline Methods

BERTOPIC (Grootendorst, 2022). To run BERTOPIC¹², we feed in the same OpenAI embeddings that we use for HYPOTHESAES, and otherwise use default choices for the algorithm: UMAP for dimensionality reduction, HDBSCAN for clustering, and c-TF-IDF to compute the top words associated with each topic. We tune the minimum cluster size hyperparameter in BERTOPIC, which controls topic granularity. The default value is 10, which often produces topics which are too granular (e.g., relative to the underlying granularity of the subtopics on WIKI and BILLS). We test all values in $\{10, 20, 50, 100, 200, 500\}$. We choose this parameter by maximizing the validation performance of the L_1 -regularized predictor which uses H features; this is identical to how we choose hyperparameters for HYPOTHESAES. We use default values of all other hyperparameters, which is standard in prior work (e.g., Pham et al. (2024)). After running BERTOPIC, we label topics using the prompt Figure 7.

NLPARAM (Zhong et al., 2024). We run NLPARAM using the publicly available implementation¹³. We set the number of candidate labels per bottleneck layer neuron to 3 (the same as our method), and we use the same text embeddings as our method (OpenAI’s text-3-embedding-small). Zhong et al. (2024) release two different prompts, one for “simple” hypotheses and one for “complex” hypotheses; following their work, we use the “simple” prompt for our 3 synthetic tasks and the “complex” prompt for our 3 real-world tasks. For all other hyperparameters, we use default values.

HYPOGENIC (Zhou et al., 2024). We run HYPOGENIC using the publicly available implementation¹⁴. We use the default parameters as defined in examples/generation.py. Following the format in their example tasks, we write custom prompts in a consistent format describing each our of tasks. The method uses the same LLM for hypothesis generation and hypothesis scoring, and we run all experiments using GPT-4o-mini (the default model in their code repository, and also the model that we use for scoring hypotheses in HYPOTHESAES).

E. Proofs

Proof of Proposition 3.1. Set $p_{ab} := \Pr[\hat{Z} = a, Z = b]$ and $y_{ab} := \mathbb{E}[Y|\hat{Z} = a, Z = b]$ for $a, b \in \{0, 1\}$. Then

$$\mathbb{E}[Y|\hat{Z} = 1] = \frac{y_{11}p_{11} + y_{10}p_{10}}{p_{11} + p_{10}} \quad (8)$$

$$\mathbb{E}[Y|Z = 1] = \frac{y_{11}p_{11} + y_{01}p_{01}}{p_{11} + p_{01}}. \quad (9)$$

Therefore,

$$\mathbb{E}[Y|\hat{Z} = 1] - \mathbb{E}[Y|Z = 1] = \frac{y_{11}p_{11} + y_{10}p_{10}}{p_{11} + p_{10}} - \frac{y_{11}p_{11} + y_{01}p_{01}}{p_{11} + p_{01}} \quad (10)$$

$$= y_{11} \left(\frac{p_{11}}{p_{11} + p_{10}} - \frac{p_{11}}{p_{11} + p_{01}} \right) + \frac{y_{10}p_{10}}{p_{11} + p_{10}} - \frac{y_{01}p_{01}}{p_{11} + p_{01}} \quad (11)$$

$$= y_{11} \left(\frac{p_{01}}{p_{11} + p_{01}} - \frac{p_{10}}{p_{11} + p_{10}} \right) + \frac{y_{10}p_{10}}{p_{11} + p_{10}} - \frac{y_{01}p_{01}}{p_{11} + p_{01}} \quad (12)$$

Since $y_{10} \leq 1$ and $y_{01} \geq 0$, (12) is bounded from above by

$$y_{11} \left(\frac{p_{01}}{p_{11} + p_{01}} - \frac{p_{10}}{p_{11} + p_{10}} \right) + \frac{p_{10}}{p_{11} + p_{10}}. \quad (13)$$

If $\frac{p_{01}}{p_{11} + p_{01}} - \frac{p_{10}}{p_{11} + p_{10}} \geq 0$, (13) is bounded from above by

$$\left(\frac{p_{01}}{p_{11} + p_{01}} - \frac{p_{10}}{p_{11} + p_{10}} \right) + \frac{p_{10}}{p_{11} + p_{10}} = \frac{p_{01}}{p_{11} + p_{01}}, \quad (14)$$

¹²<https://github.com/MaartenGr/BERTopic>

¹³<https://github.com/ruiqi-zhong/nlparam>

¹⁴<https://github.com/ChicagoHAI/hypothesis-generation>

since $y_{11} \leq 1$. Otherwise, (13) is bounded from above by

$$\frac{p_{10}}{p_{11} + p_{10}}. \quad (15)$$

Also note that (12) is bounded from below by

$$y_{11} \left(\frac{p_{11}}{p_{11} + p_{10}} - \frac{p_{11}}{p_{11} + p_{01}} \right) - \frac{p_{01}}{p_{11} + p_{10}} = y_{11} \left(\frac{p_{01}}{p_{11} + p_{01}} - \frac{p_{10}}{p_{11} + p_{10}} \right) - \frac{p_{01}}{p_{11} + p_{10}}. \quad (16)$$

If $\frac{p_{01}}{p_{11} + p_{01}} - \frac{p_{10}}{p_{11} + p_{10}} \leq 0$, (16) is bounded from below by

$$\left(\frac{p_{01}}{p_{11} + p_{01}} - \frac{p_{10}}{p_{11} + p_{10}} \right) - \frac{p_{01}}{p_{11} + p_{10}} = -\frac{p_{10}}{p_{11} + p_{10}}, \quad (17)$$

since $y_{11} \leq 1$. Otherwise, (16) is bounded from below by

$$-\frac{p_{01}}{p_{11} + p_{01}}. \quad (18)$$

Therefore,

$$|\mathbb{E}[Y|\hat{Z} = 1] - \mathbb{E}[Y|Z = 1]| \leq \max \left\{ \frac{p_{10}}{p_{11} + p_{10}}, \frac{p_{01}}{p_{11} + p_{01}} \right\}. \quad (19)$$

Analogously, we can bound $|\mathbb{E}[Y|\hat{Z} = 0] - \mathbb{E}[Y|Z = 0]|$ by

$$\max \left\{ \frac{p_{10}}{p_{00} + p_{10}}, \frac{p_{01}}{p_{00} + p_{01}} \right\}. \quad (20)$$

Notice that

$$\frac{p_{10}}{p_{00} + p_{10}} = \frac{p_{10}}{p_{11} + p_{10}} \cdot \frac{p_{11} + p_{10}}{p_{00} + p_{10}} \quad (21)$$

$$= \frac{p_{10}}{p_{11} + p_{10}} \cdot \frac{\Pr[\hat{Z} = 1]}{\Pr[Z = 0]} \quad (22)$$

Similarly,

$$\frac{p_{01}}{p_{00} + p_{01}} = \frac{p_{01}}{p_{11} + p_{01}} \cdot \frac{p_{11} + p_{01}}{p_{00} + p_{01}} \quad (23)$$

$$= \frac{p_{01}}{p_{11} + p_{01}} \cdot \frac{\Pr[Z = 1]}{\Pr[\hat{Z} = 0]} \quad (24)$$

Set $p := \min\{\Pr[\hat{Z} = 0], \Pr[Z = 0]\}$. Then $\Pr[\hat{Z} = 0], \Pr[Z = 0] \geq p$ and $\Pr[\hat{Z} = 1], \Pr[Z = 1] \leq 1 - p$

$$\frac{\Pr[\hat{Z} = 1]}{\Pr[Z = 0]}, \frac{\Pr[Z = 1]}{\Pr[\hat{Z} = 0]} \leq \frac{1 - p}{p}. \quad (25)$$

It follows that

$$|S(\hat{Z}) - S(Z)| = |(\mathbb{E}[Y|\hat{Z} = 1] - \mathbb{E}[Y|\hat{Z} = 0]) - (\mathbb{E}[Y|Z = 1] - \mathbb{E}[Y|Z = 0])| \quad (26)$$

is at most

$$\left(1 + \frac{1 - p}{p} \right) \max \left\{ \frac{p_{10}}{p_{11} + p_{10}}, \frac{p_{01}}{p_{11} + p_{01}} \right\} = \frac{1}{p} \max \left\{ \frac{p_{10}}{p_{11} + p_{10}}, \frac{p_{01}}{p_{11} + p_{01}} \right\}. \quad (27)$$

The result follows by noting that $p := \min\{\Pr[\hat{Z} = 0], \Pr[Z = 0]\}$ and $\frac{p_{10}}{p_{11} + p_{10}}, \frac{p_{01}}{p_{11} + p_{01}} = 1 - \text{recall}, 1 - \text{precision}$. $\square$

F. Prompts

You are a machine learning researcher who has trained a neural network on a text dataset. You are trying to understand what text features cause a specific neuron in the neural network to fire.

You are given two sets of SAMPLES: POSITIVE SAMPLES that strongly activate the neuron, and NEGATIVE SAMPLES from the same distribution that do not activate the neuron. Your goal is to identify a feature that is present in the positive samples but absent in the negative samples.

Example features could be:

- “uses multiple adjectives to describe colors”
- “describes a patient experiencing seizures or epilepsy”
- “contains multiple single-digit numbers”

POSITIVE SAMPLES:

{positive_samples}

NEGATIVE SAMPLES:

{negative_samples}

Rules about the feature you identify:

- The feature should be objective, focusing on concrete attributes rather than abstract concepts.
- The feature should be present in the positive samples and absent in the negative samples. Do not output a generic feature which also appears in negative samples.
- The feature should be as specific as possible, while still applying to all of the positive samples. For example, if all of the positive samples mention Golden or Labrador retrievers, then the feature should be “mentions retriever dogs”, not “mentions dogs” or “mentions Golden retrievers”.

Do not output anything else. Your response should be formatted exactly as shown in the examples above. Please suggest exactly one description, starting with “-” and surrounded by quotes “”. Your response is:

_-“

Figure 4. The prompt we use to interpret SAE neurons. The positive samples are texts that strongly activate the neuron, while the negative samples are texts that weakly activate the neuron. The “example features” section can be edited to produce features in a different format.

Is text_a and text_b similar in meaning? Respond with yes, related, or no.

Here are a few examples.

Example 1:

text_a: has a topic of protecting the environment

text_b: has a topic of environmental protection and sustainability

output: yes

Example 2:

text_a: has a language of German

text_b: has a language of Deutsch

output: yes

Example 3:

text_a: has a topic of the relation between political figures

text_b: has a topic of international diplomacy

output: related

Example 4:

text_a: has a topic of the sports

text_b: has a topic of sports team recruiting new members

output: related

Example 5:

text_a: has a named language of Korean

text_b: uses archaic and poetic diction

output: no

Example 6:

text_a: has a named language of Korean

text_b: has a named language of Japanese

output: no

Example 7:

text_a: describes an important 20th century historical event

text_b: describes a 20th century European politician

output: no

Target:

text_a: {text_a}

text_b: {text_b}

output:

Figure 5. The prompt we use to evaluate surface similarity between the inferred hypotheses and reference topics. The prompt is lightly edited from [Zhong et al. \(2024\)](#). The few-shot examples demonstrate three levels of similarity: equivalent meaning (yes), related concepts (related), and unrelated concepts (no).

Check whether the TEXT satisfies a PROPERTY. Respond with Yes or No with an explanation that discusses the evidence from the TEXT (at most a sentence). When uncertain, output No.

Example 1:

PROPERTY: "mentions a natural scene."

TEXT: "I love the way the sun sets in the evening."

Output: Yes. "Sun sets" are clearly natural scenes.

Example 2:

PROPERTY: "writes in a 1st person perspective."

TEXT: "Jacob is smart."

Output: No. This text is written in a 3rd person perspective.

Example 3:

PROPERTY: "is better than group B."

TEXT: "I also need to buy a chair."

Output: No. It is unclear what the PROPERTY means (e.g., what does group B mean?) and doesn't seem related to the text.

Example 4:

PROPERTY: "mentions that the breakfast is good on the airline."

TEXT: "The airline staff was really nice! Enjoyable flight."

Output: No. Although the text appreciates the flight experience, it DOES NOT mention about the breakfast.

Example 5:

PROPERTY: "appreciates the writing style of the author."

TEXT: "The paper absolutely sucks because its underlying logic is wrong. However, the presentation of the paper is clear and the use of language is really impressive."

Output: Yes. Although the text dislikes the paper, it says "the presentation of the paper is clear", so it DOES like the writing style.

Example 6:

PROPERTY: "has a formal style; specifically, the language in the text is relatively formal, complex and academic. For example, 'represent whom and which'"

TEXT: "investigates formation of nominalization"

Output: Yes. "formation" and "nominalization" are abstract and complex nouns.

Example 7:

PROPERTY: "refers to historical dates; specifically, there are references to years or specific dates in the text. For example, 'Obama was born on August 4, 1961.'"

TEXT: "A member of the Democratic Party, he was the first African-American president of the United States."

Output: No. The text does not mention date.

Now complete the following example - Respond with Yes or No with an explanation that discusses the evidence from the TEXT. When uncertain, output No.

PROPERTY: {hypothesis}

TEXT: {text}

Output:

Figure 6. The prompt we use to compute hypothesis annotations with GPT-4o-mini to evaluate hypothesis quality. The prompt is lightly edited from Zhong et al. (2024).

You are a machine learning researcher analyzing clusters from a topic model. You need to identify the key distinguishing features of documents in a specific topic cluster.

Example features could be:

- “uses multiple adjectives to describe colors”
- “describes a patient experiencing seizures or epilepsy”
- “contains multiple single-digit numbers”

These are the top keywords describing the topic: {topic_keywords}.

You will also see a sample of documents belonging to this topic and a random sample of documents from other topics from the same dataset.

TOPIC DOCUMENTS (a sample of documents belonging to this topic):

{topic_documents}

RANDOM DOCUMENTS (a sample of documents from other topics):

{random_documents}

Your goal is to create a short, specific label that captures what makes the topic documents different from the contrasting documents.

The label should be:

1. As specific as possible, while still applying to all topic documents
2. Focused on concrete features rather than abstract concepts
3. Based on characteristics present in the topic documents but not in the contrasting documents

Do not output anything else. Your response should be a single topic label starting with “-” and surrounded by quotes “”. Your response is:

_-“

Figure 7. The prompt we use to generate interpretable topic labels for BERTOPIC. The LLM is prompted with the keywords that distinguish the topic from others, and also with two sets of documents: 10 documents belonging to the target topic, and a random sample of 10 documents from other topics. This labeling approach is inspired by Grootendorst (2024), but we designed the prompt ourselves (intending it to be similar, where possible, to our neuron interpretation prompt).