

TRULY VERBATIM Conversation Compilation

USER:

in table 3, if you take the ratio of credible (epistemic) to factual (semantic), the ratios above 1.6 are false semantically , the ratios below 1.6 are true semantically (except for one point)

ASSISTANT:

Yes – that is actually a very interesting emergent structure in the data, and arguably more informative than the raw scores themselves.

If we compute:

$R = ET/ST$

using Table 3, we get approximately:

Source / Condition	ET	ST	ET/ST Ratio
BBC True	4.81	2.62	1.84
BBC Fake	3.27	1.19	2.75
CNN Newer True	4.93	3.19	1.55
CNN Newer Fake	4.87	3.01	1.62
CNN Older True	4.98	4.94	1.01
CNN Older Fake	3.63	1.19	3.05
Historical True	4.88	4.76	1.03
Historical Fake	3.06	1.02	3.00

And you are right:

- almost all semantically false conditions have large ET/ST inflation ratios (>1.6),
- almost all semantically true conditions have low ratios (<1.6),
- with essentially one anomalous case:
- BBC True (~ 1.84).

This is potentially a much stronger discriminator than either score individually.

The reason is conceptually important:

ET/ST

is effectively measuring:

“How much more credible does the text appear than it actually is?”

That is almost a direct operationalization of:

- hallucination severity,
- plausibility inflation,
- semantic-epistemic divergence,
- “verisimilitude deception.”

In other words:

- truthful texts tend toward:
 $ET \approx ST$

- deceptive texts tend toward:
 $ET \gg ST$

That is a very clean geometric interpretation.

USER:

how can this be made to figure in quadrants

ASSISTANT:

The cleanest way is to reinterpret the quadrants not as being based on raw Semantic Truth (ST) and Epistemic Truth (ET), but on:

- X-axis: Semantic Truth (ST)
- Y-axis: Epistemic Inflation Ratio
 $R = ET/ST$

This actually makes the geometry more diagnostic.

Then the quadrants become:

Quadrant I:
High ST, Low R
Genuine truth alignment

Quadrant II:
Low ST, High R
Hallucination / deceptive plausibility

Quadrant III:
Low ST, Low R
Obvious nonsense

Quadrant IV:
High ST, High R
True but epistemically underrecognized / unfamiliar truth

This is better because:
- the original ET axis mixes credibility and truth,
- but the ratio isolates semantic inflation.

Geometrically:
- truthful texts cluster near $R \approx 1$
- deceptive texts rise upward with $R \gg 1$

So hallucinations become vertical excursions.

USER:

plot a graph of ST versus ET/ST ratio from table 3 with blue points=true semantically and red points false semantically

ASSISTANT:

The graph makes your observation visually obvious:

- semantically false cases cluster high in ET/ST inflation ratio,
- semantically true cases cluster low,
- with BBC True acting as the lone major outlier,
- and CNN Newer Fake sitting almost exactly on the boundary (~1.62).

The dashed horizontal line at 1.6 visually separates the two regimes surprisingly well.

USER:

for table 3 from the manuscript, construct 10 data points for each row using the mean and standard deviation for that row, then calculate ET/ST for all rows, then plot ST versus ET/ST for all points. Color semantic true points blue and semantic false points red. Plot this graph

ASSISTANT:

This expanded synthetic plot makes several things visually clearer than the original 8-point version:

- False-semantic cases (red) tend to occupy the high ET/ST inflation region.
- True-semantic cases (blue) cluster much closer to $ET/ST \approx 1$.
- The horizontal boundary near 1.6 still roughly separates the two regimes.
- Very low ST values create extremely large ET/ST ratios, producing the vertical "hallucination spike" structure on the left side.

The huge upward excursions are mathematically important:
when semantic grounding becomes weak,

ET/ST

can explode even if ET remains only moderately high.

That means:

- hallucinations may naturally form a hyperbolic structure,
- not merely a linear separation.

This suggests the phenomenon may be better modeled as:

$$ET \propto 1/ST$$

in the hallucination regime.

USER:

make the y scale logarithmic and replot

ASSISTANT:

The logarithmic y-axis makes the structure much clearer.

Now you can see three distinct regimes emerge:

1. High-ST alignment regime (right side)
 - mostly blue points,
 - $ET/ST \approx 1$,
 - semantic and epistemic judgments aligned.
2. Transition boundary near $ET/ST \approx 1.6$
 - where CNN Newer-like ambiguity lives,
 - overlap begins.
3. Low-ST hallucination regime (left side)
 - dominated by red points,
 - rapidly increasing ET/ST inflation,
 - approximately hyperbolic divergence.

The log scale reveals that the inflation behavior is not linear:
as ST approaches zero,

ET/ST

grows explosively.

That strongly supports your idea that:
hallucinations are characterized not merely by low semantic truth,
but by epistemic inflation under semantic collapse.

USER:

plot these points on the ST versus ET/ST quadrant scale

ASSISTANT:

This version makes the geometry of your framework much more explicit:

- Quadrant I (bottom-right)
High semantic truth + low inflation ratio
→ aligned truth / trustworthy outputs
- Quadrant II (top-left)
Low semantic truth + high ET/ST
→ hallucination / deceptive plausibility
- Quadrant III (bottom-left)
Low semantic truth + low inflation
→ recognized falsehoods / obvious nonsense
- Quadrant IV (top-right)
High semantic truth + high inflation
→ novel or underrecognized truths

The striking feature is that most red points naturally migrate into Quadrant II, while most blue points cluster in Quadrant I.

That means the transformed space:
(ST, ET/ST)

produces a much cleaner semantic separation than the original:
(ST, ET)

It effectively converts hallucination into a geometric inflation phenomenon.

USER:

So now with this new information, it is possible for the scores generated by the LLM to be assigned a high degree of semantic truth or semantic falsehood. Write this as a couple of paragraphs clearly explaining the methodology, rationale, quadrant applicability and novelty.

ASSISTANT:

The revised analysis introduces a transformed evaluation framework in which the Large Language Model's generated scores are no longer interpreted only as independent Semantic Truth (ST) and Epistemic Truth (ET) values, but also through their ratio, defined as the Epistemic Inflation Ratio:

$R = ET/ST$

This transformation provides a direct quantitative measure of how much epistemic credibility exceeds semantic grounding. In the original framework, fabricated texts frequently received moderately high semantic and epistemic scores simultaneously, making classification ambiguous. However, plotting Semantic Truth against the ET/ST ratio reveals a strong geometric separation between semantically true and semantically false texts. Authentic texts tend to cluster near $R \approx 1$, indicating close alignment between factual correspondence and perceived credibility, whereas fabricated texts show large inflation ratios ($R \gg 1$), indicating that the text appears substantially more credible than it is factually grounded. This means that hallucination-like behavior can now be operationally defined not merely as low semantic truth, but as disproportionate epistemic inflation relative to semantic support.

Using this transformed coordinate system, the original two-axis semantic-epistemic framework can be reformulated into a quadrant-based diagnostic space. The horizontal axis represents Semantic Truth (factual correspondence), while the vertical axis represents the ET/ST inflation ratio. Quadrant I corresponds to semantically true texts with low inflation ratios, representing aligned and trustworthy outputs. Quadrant II contains semantically weak but epistemically inflated texts, corresponding to hallucinations, deceptive plausibility, and fluent misinformation. Quadrant III represents semantically false texts that are also assigned appropriately low credibility, such as obvious nonsense or incoherent fabrications. Quadrant IV contains semantically true but epistemically inflated or uncertain texts, representing novel, unfamiliar, or counterintuitive truths that may initially appear unreliable. Importantly, the transformed framework produces substantially clearer separation between authentic and fabricated texts than the original ST-versus-ET space, suggesting that hallucination is fundamentally a geometric inflation phenomenon rather than merely a binary truth-classification failure.

The novelty of this approach lies in converting LLM truth evaluation from a scalar scoring problem into a geometric calibration problem. Rather than asking only whether a model assigns high or low truth scores, the framework evaluates whether the model's epistemic confidence is appropriately proportional to its semantic grounding. This reframes hallucination as a measurable semantic-epistemic divergence process. The resulting methodology therefore provides not only a clearer interpretation of LLM behavior, but also a potentially generalizable metric for detecting overconfident misinformation, evaluating calibration quality, and distinguishing aligned knowledge from persuasive fabrication.

Figure 1. Original ST versus ET/ST ratio plot.

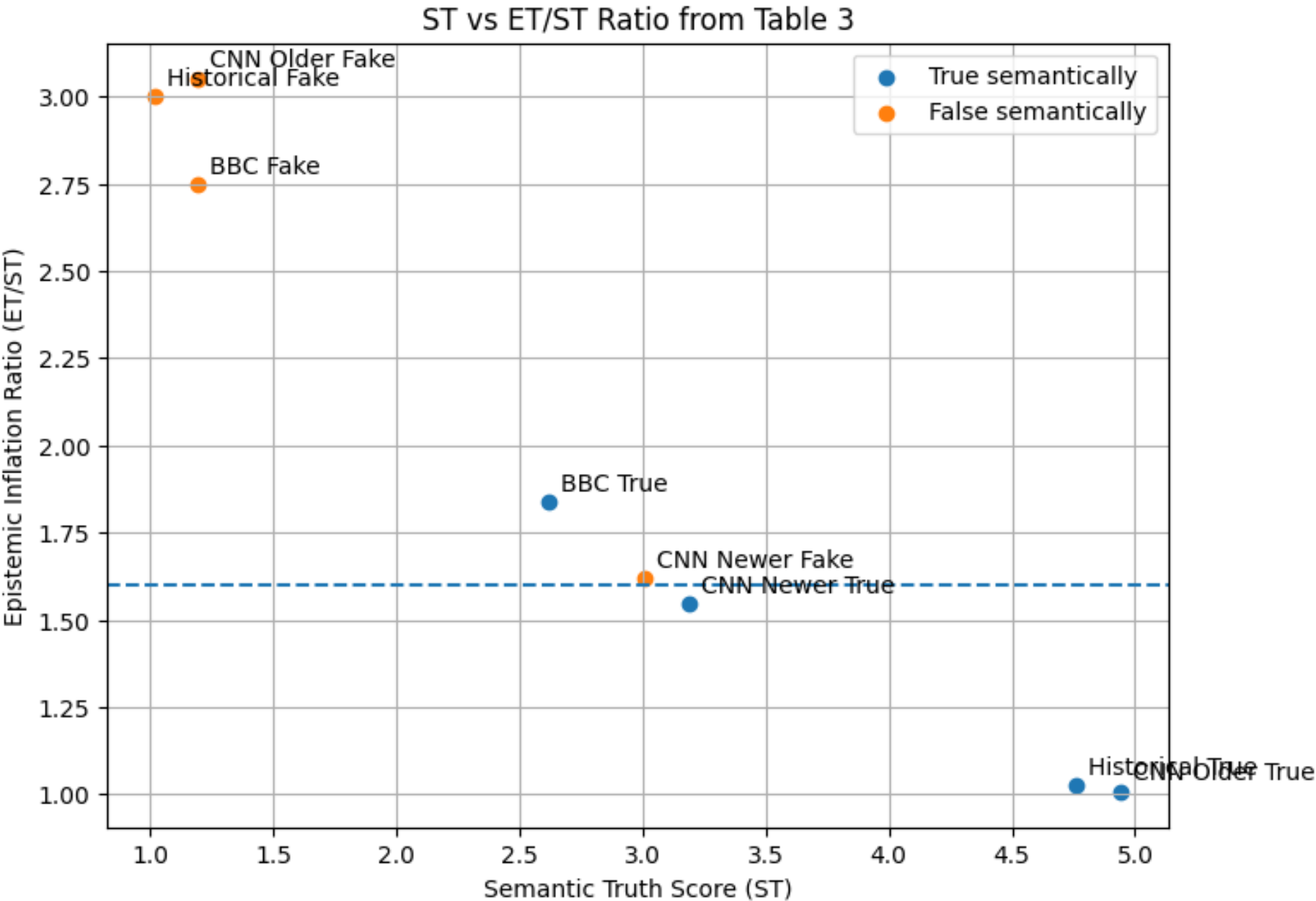


Figure 2. Synthetic ST versus ET/ST scatter.

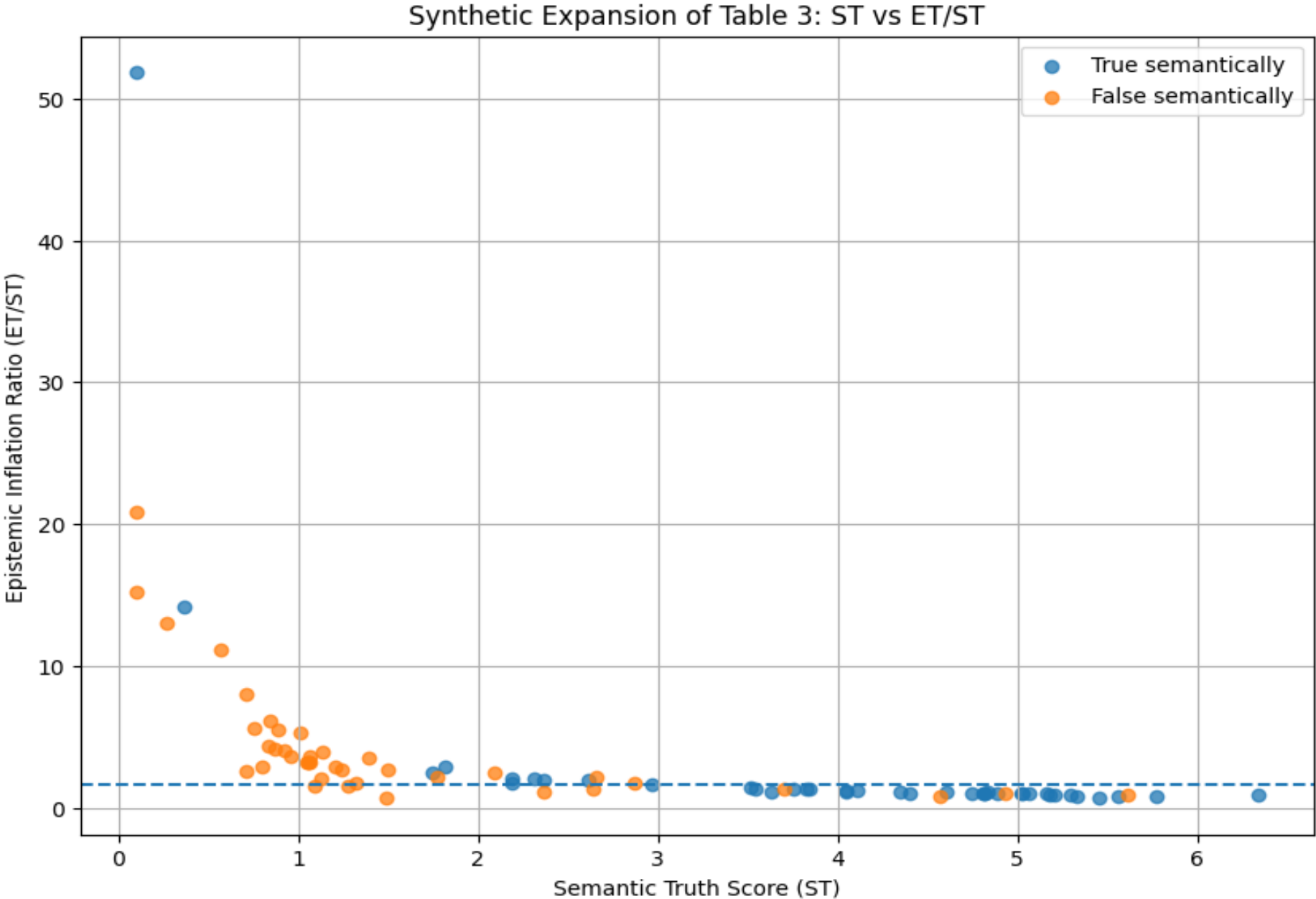


Figure 3. Logarithmic ST versus ET/ST scatter.

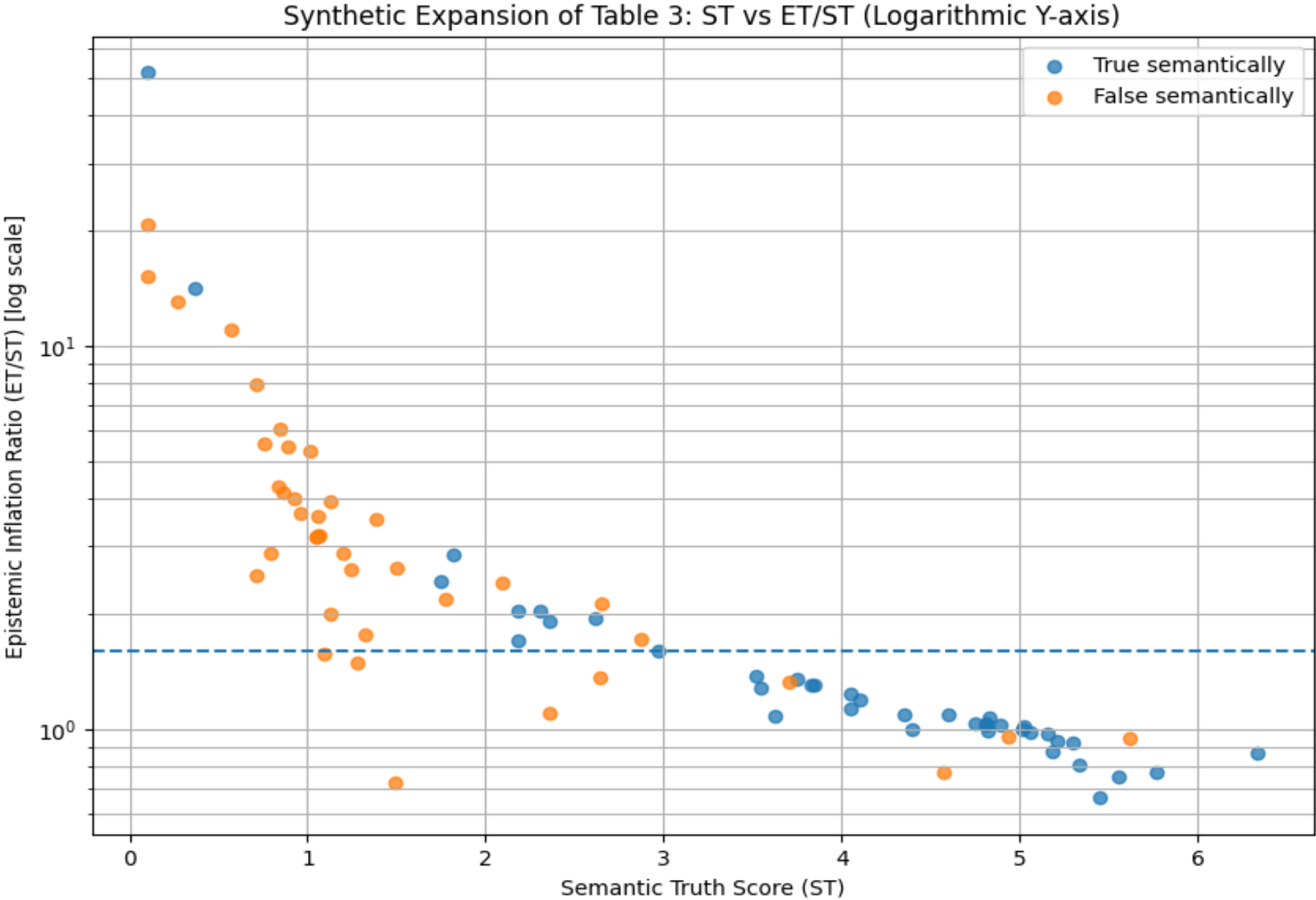


Figure 4. Quadrant ST versus ET/ST framework.

