

プレスリリース

2026 年 9 月 14 日

有限会社オーセンエンジニアリング
SAS 研究所 (SEMARCHST LAB)

AI へ送る情報量を 81.8%削減、評価対象 100 意味を 100%保持

独自開発の意味選択 AI エンジン「SSAE」で実証

— 意味を選んで AI に考えさせる「省計算」で、クラウド AI からローカル LLM まで —

有限会社オーセンエンジニアリング(山形県鶴岡市)は、弊社 SAS 研究所で独自に研究開発を進めている「SSAE (Selective Semantic AI Engine / 意味選択 AI エンジン)」の基礎実証において、AI (LLM) へ送る情報量を、本実験条件における通常 LLM 相当の比較基準に対して『81.8%削減』しながら、評価対象とした『100 意味を 100%保持 (100/100 PASS)』する結果を確認しました。SSAE は、LLM へ渡す前に現在の処理に必要な意味を選択し、AI に考えさせる意味の範囲そのものを絞る外部意味処理エンジンです。今回の成果は、AI そのものを高性能化するだけでなく、『AI の能力を必要なところだけに使う』という新たな AI 処理効率化の可能性を示すものです。

『意味構築学』から『SAS OS』、そして『SSAE』へ

SSAE は、単独のプロンプト圧縮技術として開発されたものではありません。弊社が研究する『意味構築学』を理論的基盤とし、意味を LLM の外部で保持・管理し、意味の境界や関係を扱う『SAS OS』の研究開発を経て、その時の処理に必要な意味・条件・関係を選択して LLM へ渡すエンジンとして開発したものです。文章や token を削ること自体を目的とするのではなく、『何を残すべきかを意味から判断し、AI に考えさせる範囲そのものを制御する』ことに特徴があります。

81.8%削減、100/100 PASS

今回の実験では、本実験系列で通常 LLM 相当の比較基準として設定した Trim Baseline 約 251.0 input tokens に対し、SSAE による入力は平均 45.8 input tokens となり、『約 81.8%削減』しました。同時に、評価対象 100 意味について『100/100 PASS、意味 FAIL 0』を確認しました。重要なのは、単に情報量を減らしたことなく、『必要な意味を保持したまま、AI が処理する情報範囲を大幅に小さくした』ことです。意味構築学から SAS OS、SSAE へと積み上げてきた研究が、具体的な数値として確認された成果です。

この成果がつながる先

生成 AI の急速な普及に伴い、計算資源、電力、データセンターへの需要が増大しています。SSAE によって一回一回の AI 処理から不要な情報処理を減らすことができれば、同じ計算資源をより有効に利用し、クラウド AI では電力需要やデータセンター設備増強への圧力を需要側から緩和する可能性があります。一方、今後普及が見込まれる『ローカル LLM』では、PC、企業内サーバー、工場設備、ロボットなど限られた計算環境で AI を動かすため、処理効率がより重要になります。SSAE は『AI に考えさせる意味の範囲を必要な部分に絞る』ことで、限られた計算資源を有効に使い、ローカル LLM の実用性を高める可能性があります。

今後の展開・9 月 27 日に公開実証

今後は、異なる AI モデルでの再現性、実計算量、処理時間、API コスト、電力消費、CO₂排出量との関係を検証し、ローカル LLM、企業 AI、製造、ロボット、行政、インフラなどへの応用を進めます。なお、今回の 81.8%の入力情報量削減が、そのまま同率の電力・CO₂削減を意味するものではありません。2026 年 9 月 27 日開催の『つるおか環境フェア』では、SSAE と通常 AI を比較し、AI の処理効率と環境負荷を可視化する実証展示を予定しています。

『AI にも、必要なところだけ考える仕組みを。』

SSAE は、LLM に考えさせる意味の範囲を必要な部分に絞り、『無駄な計算を減らし、限られた計算資源をより有効に使う』。そして『増大する AI 需要を、より少ない計算資源で支える』技術を目指します。