

gpt2

124M parameters · 12 layers · scanned 2026-07-24T16:43:25Z



GRADE C · FAIR

K# 92.8

100 = BASELINE ACROSS THE INDEX

02 Executive brief
06 Architecture
09 Compression density
12 Capability
18 About

04 Executive summary
07 Activation profile
10 Performance
14 Knockout analysis

05 Peer comparison
08 Layer analysis
11 Coherence & efficiency
16 Prescriptions

Confidential · prepared for the model owner · Keel Technologies, Inc.

What this report says.

Your model scores **K# 92.8** on the Keel Index — a below-baseline result, grade C. It ranks **1 of 2** in the Micro class. Your model tops its size class.

Three findings.

34%

of compute is not doing
useful work

3 of 12

layers can be removed with
no measured quality loss

**\$1,418–
\$4,760**

per month in avoidable
spend at 1B tokens/month

Three actions, in order of impact.

1. Fine-tune on domain data (5–15% quality gain on-task is typical, without touching structure)
2. Consider LoRA adapters (add task-specific capacity at ~1% parameter cost)
3. Evaluate the class leader or a modern small-model alternative if size is not a hard constraint

Next step. A **Keel Forge** engagement can scope the fine-tuning or model-swap work on your own eval set. Contact contact@keeltech.co to scope.

One diagnosis, five readings.

This report looks inside gpt2 — not at what it says, but at how it is built. Each reading below is a lens on a different property of the model's structure.

Keel Score (K#)	A single number for structural quality, positioned against every model we have scanned.
------------------------	---

Architecture Class	The structural shape the model falls into, and what it predicts about behavior.
---------------------------	---

Coherence	How evenly capability is distributed across the model rather than concentrated.
------------------	---

Thermal Grade	Efficiency — how much compute does useful work versus sits idle.
----------------------	--

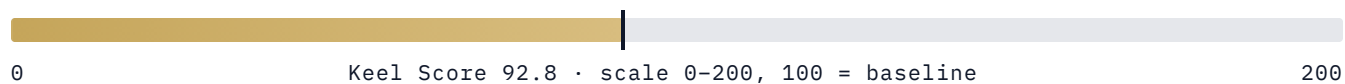
Knockout analysis	Which layers can be removed with little measurable impact.
--------------------------	--

The measurement uses a proprietary methodology. This report shows the result, not the method. No training data was required and nothing about the model leaves the engagement.

EXECUTIVE SUMMARY

gpt2 scores 92.8, grade C.

Against models of comparable size (Micro (< 200M)), gpt2 is a **fair** result. Its structural shape is **Declining** — **capability tapers gradually across depth**, and efficiency is graded **C** — 66% of its compute is doing useful work.



AT A GLANCE

READING	RESULT	RATING
Keel Score (K#)	92.8	C Fair
Coherence Score	79.6 / 100	Good
Architecture Class	Declining	Tapers with depth
Thermal Grade	C	Fair
Perplexity (held-out)	19.52	Fair for size
Redundant layers (sampled)	3 of 6	Optimization headroom

Verdict. A middling structural result in the Micro (< 200M) class. The knockout analysis identified layers 3, 6, 8 as redundant at no measured cost — the fastest path to lower inference cost. The thermal grade of C means roughly 66% of compute does useful work; there is real room to reclaim.

PEER COMPARISON

Rank 1 of 2 in your size class.

gpt2 sits in the **Micro (< 200M)** band on the Keel Index. Here is where your model ranks against every other scanned model in the same band.

YOUR MODEL

gpt2 ← your model

K# 92.8

C

Your model tops its size class. No scanned peer in the Micro (< 200M) band scored higher.

BELOW YOUR MODEL

distilgpt2

K# 83.0

C

The Keel Index is growing continuously. Your ranking reflects models scanned to date and will refine as the database expands.

Declining.

Every model resolves to one of five structural classes. gpt2 is **Declining**: its capability tapers gradually across depth, with early layers doing more work than later ones. Declining models often have architectural room in the tail — later layers may be candidates for pruning.

PARAMETERS 124M	LAYERS 12	WIDTH 768
CONTEXT 1024	CLASS Declining	UTILIZATION 66%

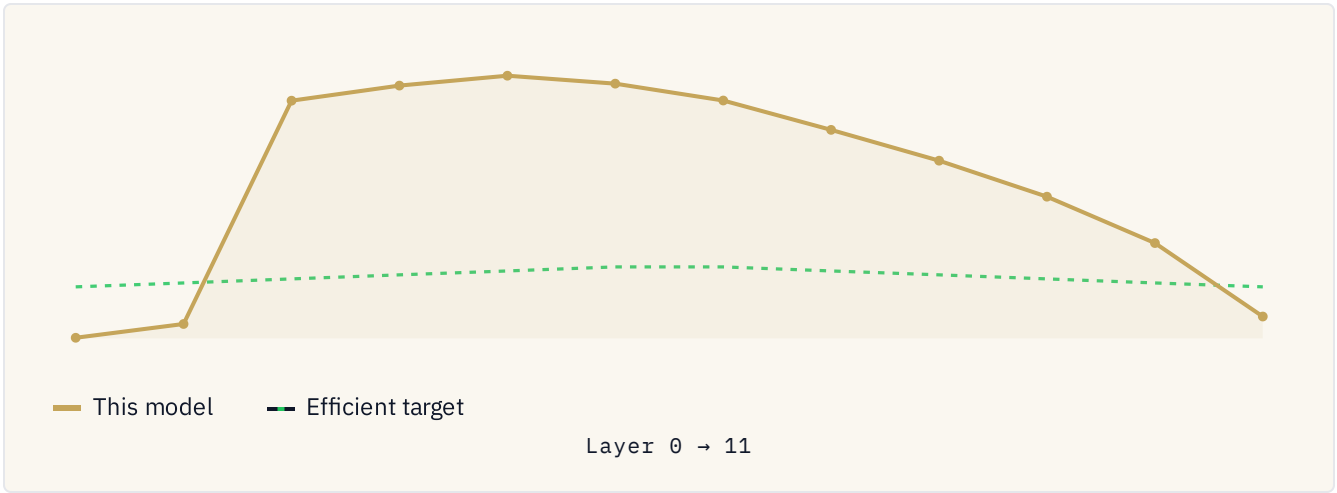
THE FIVE CLASSES

Sequential	Concentrated	Front-Loaded	Declining	Dual-Phase
------------	--------------	--------------	------------------	------------

Class is a structural signature, not a quality grade. It tells you how the model will behave — under fine-tuning, under pruning, under distribution shift — more than how good it is.

The gap is the opportunity.

The solid line is gpt2’s activation variance across its 12 layers. The dotted line below it is the profile of an efficient, evenly-built model at this size. The distance between them is exactly what an optimization pass would recover.



PEAK VARIANCE layer 4	QUIETEST LAYER layer 0	SHAPE Declining
--------------------------	---------------------------	--------------------

Where the work happens.

Activation variance for each of the 12 layers. Variance tapering across depth is the Declining signature — later layers do progressively less.

LAYER	ACTIVATION VARIANCE		LEVEL
L0	0.7		Low
L1	13.1		Low
L2	215.6		Peak
L3	229.3		Peak
L4	238.3		Peak
L5	231.1		Peak
L6	215.7		Peak
L7	189.2		High
L8	161.2		High
L9	128.5		High
L10	86.5		Moderate
L11	19.8		Low

Bits per parameter, calibrated against the Index.

Compression Density measures how much structural information the model stores per parameter, calibrated against every model on the Keel Index. This reading sits alongside K# and Grade — it measures an aspect of structure the headline number aggregates away, and it is computed directly from the model's weights, with no forward passes or training data required.

READING

86% of baseline

Slightly below

Above baseline means the model is more parameter-efficient than average for its size class — it stores more structural information per parameter than the Index median.

Below baseline means the opposite: there is structural room for the same capability at lower parameter count.

How this relates to the K#. Two models with the same K# can have very different Compression Density readings. A high-density model is punching above its weight class; a low-density model is achieving the same result with parameter redundancy. Together with Knockout Analysis on the next-but-one page, this reading tells the customer whether their model's structural room is at the layer level (Knockout) or the parameter-count level (Compression Density).

How well it predicts.

Standard held-out metrics, contextualized for a model. These measure output quality; the rest of the report measures structure.

Perplexity (held-out)

19.52 Fair for size



lower is better

reference class:

Bits per byte

4.287 compression density



lower is better

drives the Keel Score

Perplexity is measured on held-out text the model never trained on. Training-set numbers measure memorization, not quality, and are never reported here.

Evenly built, partly idle.

COHERENCE SCORE

79.6_{/100}



Good — capability is reasonably distributed, with some concentration early.

THERMAL GRADE

C

A

B

C

D

F

Fair — roughly 66% of compute is doing useful work. The remainder is headroom.

Coherence and efficiency are independent of raw score. A high-scoring model can still be uneven or wasteful — which is precisely where the value of a scan lies.

What models at this level tend to do.

These are **correlations from our scan database**, not guarantees and not causal claims. They describe what has been observed in models at a given structural level — they do not certify that this model will or will not perform a task.

CAPABILITY	TYPICALLY SEEN AT	THIS MODEL
Byte-level compression	Baseline+	observed
Short-sequence continuation	Baseline+	observed
Local pattern recognition	Emerging+	not yet observed
Consistent formatting	Emerging+	not yet observed
Multi-step coherence	Developing+	not yet observed
Cross-domain transfer	Advanced+	not yet observed
Long-range consistency	Advanced+	not yet observed
Instruction adherence	Frontier+	not yet observed
Robustness under shift	Frontier+	not yet observed

gpt2 sits at the **Developing** tier. Capabilities marked “not yet observed” are those our database associates with higher structural levels; they are an indication of headroom, not a verdict on the model.

Reading the certificate.

The tiers below describe increasing structural sophistication. They are drawn from the models in our scan database and are updated as the database grows.

Baseline	Compression and short continuation. The foundation every model reaches.	
Emerging	Local patterns and consistent formatting begin to appear.	
Developing	Multi-step coherence starts to hold.	◀ YOUR MODEL
Advanced	Cross-domain transfer and long-range consistency.	
Frontier	Instruction adherence and robustness under shift.	

A model can be structurally excellent and still sit at a lower tier if it is small. Tier reflects observed capability range, not effort or design quality.

3 redundant layers: 3, 6, 8.

We remove layers one at a time and measure the impact on held-out quality. A value near 1.00× means the model barely notices — a candidate for pruning. Sampled across the depth of the model:

LAYER REMOVED	IMPACT ON QUALITY	VERDICT
Layer 0	145.288×	Essential
Layer 3	1.003×	Redundant
Layer 4	1.076×	Contributing
Layer 6	0.983×	Redundant
Layer 8	1.032×	Redundant
Layer 9	1.064×	Contributing

Finding. Layers 3, 6, 8 showed negligible impact when removed from the sample of 6. Pruning them is the fastest route to lower inference cost at no measured loss. A Keel Forge engagement can execute this end-to-end.

What idle capacity is costing you.

Efficiency grade C implies roughly **66%** of gpt2’s parameters are doing useful work — the remaining **34%** is compute you provision, power, and pay for on every request.



WHAT 34% IDLE COSTS AT PRODUCTION VOLUME

MONTHLY INFERENCE VOLUME	TOTAL COMPUTE COST	RECLAIMABLE WASTE
100M tokens (light production)	\$417–\$1,400/mo	\$142–\$476/mo
1B tokens (typical production)	\$4,170–\$14,000/mo	\$1,418–\$4,760/mo
10B tokens (heavy production)	\$41,700–\$140,000/mo	\$14,178–\$47,600/mo

How to read this. Ranges span typical cloud pricing for a model of this size — from RunPod-tier rates (\$1.50/hr A100) to hyperscaler-tier rates with modest reserved-capacity discounts (\$4–5/hr A100). Idle share maps close to linearly onto compute cost for transformer inference.

Actual savings depend on your serving stack, batch configuration, and traffic pattern. A Keel Forge engagement measures the before-and-after on your specific infrastructure rather than estimating it.

What to do next.

Prioritized actions drawn from this report's findings, ordered by expected impact.

HIGH

Prune layers 3, 6, 8.

DO Layers 3, 6, 8 showed $\leq 5\%$ PPL degradation when removed. Remove and briefly retrain (10–20% of original training compute) to lock in the smaller model. Parameter reduction: 25%.

EXPECT At 1B tokens/month, roughly \$1,042–\$3,500/month savings (range spans cloud pricing tiers). Same quality, fewer parameters, faster inference.

MEDIUM

Reduce weight precision (quantization).

DO Independent of structural score, applicable to every model. Fp16 $\rightarrow$ int8 or int4 typically achieves 2–4 $\times$ memory footprint reduction. Validate quality loss on your own eval set before shipping.

EXPECT Memory drops by the compression ratio; latency and cost drop proportionally on memory-bound workloads. int8 typically $\leq 1\%$ quality loss on structured tasks; int4 requires careful calibration.

LOW

Re-scan after optimization.

DO Re-scan after any change and compare against this report.

EXPECT Verified improvement on the same measurement scale.

How this was measured.

The Keel Scan reads a model's internal activation structure directly from its weights, using a proprietary methodology developed by Keel Technologies. No training data is required; the model is not modified.

Grade (A–F)	The headline result, anchored at K# = 100 (baseline). A at K# ≥ 140, B at 100–139 (above baseline), C at 75–99, D at 50–74, F below 50.
Keel Score (K#)	The underlying structural-quality score, positioned against our scan database. 100 is the baseline; above 100 is above average, below 100 is below.
Coherence	A measure of how evenly activation structure is distributed across layers.
Architecture Class	A structural signature computed from the activation profile.
Thermal Grade	An efficiency reading — the share of compute engaged in useful work.
Knockout	Empirical layer removal, measured against held-out quality.

GRADE COLOURS

 A · Excellent  B · Good  C · Fair  D · Poor  F · Failing

Held-out evaluation uses text the model never trained on. We report what we can measure and mark what we cannot; capability correlations are labeled as such.

Measurement for intelligence.

Keel is a measurement company. The Keel Scan is the first instrument — a structural diagnosis of any AI model. Comparison, monitoring, optimization, and certification follow.

Keel Scan This report. \$1,500 per model.	Keel Compass Same-model rescan — quantify what changed.
Keel Pulse Continuous monitoring and drift alerts.	Keel Forge Optimization — cut the waste this scan found.
Keel Fleet Portfolio management across many models.	Keel Certify Certification and trust mark.

Ready to act on this report? contact@keeltech.co · keeltech.co

KEEL Technologies

We grow intelligence.

Report prepared for the owner of gpt2

Scanned 2026-07-24T16:43:25Z

Keel Technologies, Inc. · Nashville, TN · keeltech.co

CONFIDENTIAL. This report and its contents are provided to the model owner under the terms of engagement. The measurement methodology is a trade secret of Keel Technologies, Inc.