

CẨM NANG CÀI ĐẶT & SỬ DỤNG AIRLLM BY LYOGAVIN

Chạy Mô Hình 70B Trên GPU 4GB VRAM · Kỹ Thuật Layer-by-Layer Inference · Tối Ưu Bộ Nhớ
Đỉnh Cao



G
i
t
H
u
b
R
e
p
o
s
i
t
o
r
y
:
<https://github.com/lyogavin>

v
i
n
/
a
i
r
l
l
m

(

-

3

3

.

8

0

0

+

S

t

a

r

s

)



T

á

c

g

i

ã

/

T

ổ

c

h

ứ

c

:

G

a

v

i

n

L

i

(

I
y
o
g
a
v
i
n
)
C
i
ấ
y
p
h
é
p
:
A
p
a
c
h
e
-
2
.
O
C
ô
n
g
n
g
h
ệ
t
h
e
n
c
h
ồ
t
:
L
a
y

e
r
-
w
i
s
e
D
i
s
k
S
t
r
e
a
m
i
n
g
,
B
l
o
c
k
-
b
y
-
b
l
o
c
k
E
x
e
c
u
t
i
o
n
,
Z
e
r

o
M
e
m
o
r
y
S
p
i
k
e
s
M
ô
h
i
n
h
h
ô
t
r
ø
:
L
l
a
m
a
3
7
0
B
,
Q
w
e
n
7
2
B
,
M
i
x
t
r

a
l
8
x
7
B
,
D
e
e
p
S
e
e
k
-
V
2
,
F
a
l
c
o
n
4
O
B
P
h
à
n
c
Ú
n
g
t
Ó
i
t
h
i
ẻ
u
:
1
C
a

r
d
đ
ồ
h
ọ
a
p
h
ổ
t
h
ô
n
g
4
C
B
V
R
A
M
(
C
T
X
1
6
5
0
,
R
T
X
3
0
5
0
.
.
.)
v
à
ổ
c
ứ
n

1. AIRLLM LÀ GÌ & PHÉP MÀU NÀO GIÚP CHẠY 70B TRÊN 4GB VRAM?

Theo cách tính thông thường, để nạp một mô hình 70 tỷ tham số (70B) ở định dạng 16-bit vào bộ nhớ, bạn cần ít nhất 140GB VRAM — tương đương 2 chiếc card đồ họa máy chủ NVIDIA A100 có giá hàng chục nghìn USD. Ngay cả khi lượng tử hóa (quantization) xuống 4-bit, bạn vẫn cần tối thiểu 35GB-40GB VRAM.

AirLLM do kỹ sư Gavin Li phát minh đã phá bỏ hoàn toàn định kiến này bằng kỹ thuật phân lớp suy luận (Layer-by-Layer Inference).

- **Nguyên lý truyền luồng theo lớp (Streaming Layers):** Mô hình Transformer thực chất gồm 80 lớp xếp chồng lên nhau. AirLLM chỉ nạp ĐÚNG 1 lớp vào VRAM tại một thời điểm, thực hiện xong phép tính thì giải phóng ngay và nạp lớp tiếp theo từ SSD vào.
- **Giữ nguyên 100% chất lượng mô hình gốc:** Không bắt buộc phải nén lượng tử hóa làm giảm độ thông minh; bạn có thể chạy mô hình ở độ chính xác đầy đủ FP16.
- **Chạy được trên cả GPU phổ thông hoặc MacBook cũ:** Card màn hình chỉ cần 4GB VRAM là đủ chỗ trống cho 1 lớp đơn lẻ tính toán.

2. YÊU CẦU HỆ THỐNG

- **GPU:** Bất kỳ card NVIDIA nào có tối thiểu 4GB VRAM hỗ trợ CUDA (GTX 1060, RTX 2060, RTX 3060...).
- **Ổ cứng (Bắt buộc):** Ổ SSD NVMe tốc độ đọc cao (khuyến nghị 3000MB/s+) và trống từ 150GB dung lượng để lưu weights mô hình.
- **RAM máy tính:** Tối thiểu 16GB RAM hệ thống.
- **Python:** Python 3.10+ cùng PyTorch phiên bản hỗ trợ CUDA.

3. HƯỚNG DẪN CÀI ĐẶT & CHẠY THỬ

Cài đặt gói AirLLM qua pip

```
r  
l  
l  
m
```

Đoạn code Python suy luận mô hình 70B (test_airlm.py)

```
f  
r  
o  
m  
a  
i  
r  
l  
l  
m  
i  
m  
p  
o  
r  
t  
A  
u  
t  
o  
M  
o  
d  
e  
l  
  
#  
K  
h  
ở  
i  
t  
ạ  
o  
m  
ô  
h  
ì  
n  
h  
L  
l  
a
```

m
a
-
3
7
0
B
(
t
ψ
đ
ộ
n
g
s
t
r
e
a
m
t
ử
n
g
l
ờ
p
v
à
o
4
G
B
V
R
A
M
)
m
o
d
e
l
=
A
u
t
o
M
o

d
e
l
.
f
r
o
m
-
p
r
e
t
r
a
i
n
e
d
(
"
m
e
t
a
-
l
l
a
m
a
/
M
e
t
a
-
L
l
a
m
a
-
3
-
7
0
B
-
I

```
n  
s  
t  
r  
u  
c  
t  
")  
  
i  
n  
p  
u  
t  
_  
t  
e  
x  
t  
=  
["  
H  
ã  
y  
g  
i  
à  
i  
t  
h  
í  
c  
h  
n  
g  
ắ  
n  
g  
o  
n  
n  
g  
u  
y  
ê  
n  
l  
ý
```

c
ủ
a
l
l
y
c
h
à
p
d
ã
n
t
h
e
o
t
h
u
y
ề
t
t
ư
ơ
n
g
đ
ồ
i
r
ộ
n
g
:
"
]

i
n
p
u
t
_
t
o
k
e
n
s

```
=  
model.  
tokenizer(  
input_text,  
return_tensors="pt",
```

```
padding = True,  
truncation = True,  
max_length = 128)  
  
# B
```

ắt đầu sử dụng ngôn ngữ generation - output put = model . generate (input

```
t
-
t
o
k
e
n
s
[
"
i
n
p
u
t
-
i
d
s
"
]
.
c
u
d
a
(
)
,
m
a
x
-
n
e
w
-
t
o
k
e
n
s
=
5
0
,
u
```

```
se-  
c-  
a-  
c-  
h-  
e-  
=  
T  
r  
u  
e,  
  
r-  
e-  
t-  
u-  
r-  
n  
-  
d-  
i-  
c-  
t  
-  
i-  
n  
-  
g-  
e-  
n-  
e-  
r-  
a-  
t-  
e-  
=  
T  
r  
u  
e)  
  
o  
u  
t  
p  
u
```

t
=
m
o
d
e
l
.
t
o
k
e
n
i
z
e
r
.
d
e
c
o
d
e
(
g
e
n
e
r
a
t
i
o
n
_
o
u
t
p
u
t
.
s
e
q
u
e
n
c

```
es[0])
print(
    output
)
```

4. LƯU Ý VỀ TỐC ĐỘ & ĐÁNH ĐỔI HIỆU NĂNG

ĐÁNH ĐỔI
I
G
I
Ữ
A
B
Ộ
N
H
Ở
V
À
T
ỐC
Đ
Ộ
V
ì
p

h
ả
i
l
i
è
n
t
ụ
c
đ
ọ
c
g
h
i
c
á
c
l
ờ
p
m
ô
h
ì
n
h
t
ử
ổ
c
ử
n
g
S
S
D
v
à
o
C
P
U
,
t
ổ
c

đ
ộ
s
i
n
h
t
o
k
e
n
s
ẽ
p
h
ụ
t
h
u
ộ
c
t
r
ự
c
t
i
ế
p
v
à
o
t
ồ
c
đ
ộ
đ
ọ
c
c
ủ
a
ồ
c
ứ
n
g

S
S
D
(
t
h
ư
ờ
n
g
đ
ạ
t
k
h
o
ả
n
g
1
-
3
t
o
k
e
n
/
g
i
â
y
t
r
ê
n
S
S
D
N
V
M
e
)
·
p
h
ư

Ơ
n
g
p
h
á
p
n
à
y
c
ự
c
k
ý
h
o
à
n
h
à
o
c
h
o
v
i
ệ
c
c
h
ạ
y
b
a
t
c
h
j
o
b
,
t
ổ
n
g
h
ợ

p
t
à
i
l
i
ệ
u
,
x
ử
l
ý
n
ề
n
h
o
ặ
c
t
h
ử
n
g
h
i
ệ
m
n
g
h
i
ê
n
c
ử
u
m
à
k
h
ô
n
g
p
h
ả

i
b
ồ
r
a
h
à
n
g
t
r
ă
m
t
r
i
ệ
u
đ
ồ
n
g
m
u
a
c
a
r
d
c
h
u
y
ê
n
d
ự
n
g
.

Chúc bạn trải nghiệm sức mạnh của những siêu mô hình 70B ngay trên cỗ máy tính cá nhân cùng AirLLM!

© Kênh AI Mỗi Ngày | Hướng dẫn AI & Linux Tự Động Hóa

Nội dung kỹ thuật tổng hợp từ tài liệu chính thức của repo <https://github.com/lyogavin/airllm>