

1. 前期准备

本文档目的是帮助在本地linux服务器端，部署并成功运行Meta开源LLM llama2。基于本文档已经成功在A100 和8卡4090服务器部署llama2-7b-chat-hf。

开始本文档，需要提前准备：

1. 支撑llama2运行的GPU服务器资源，使用基于Linux内核的服务器操作系统。
2. 确保可以通过远程连接，连接到服务器资源。（比如使用vscode remote-ssh)
3. 确保有足够的存储空间，最好有50G的富裕空间。
4. 服务器已经预装python3, CUDA等基础软件，由于本教程使用docker部署，所以不需要安装miniconda。

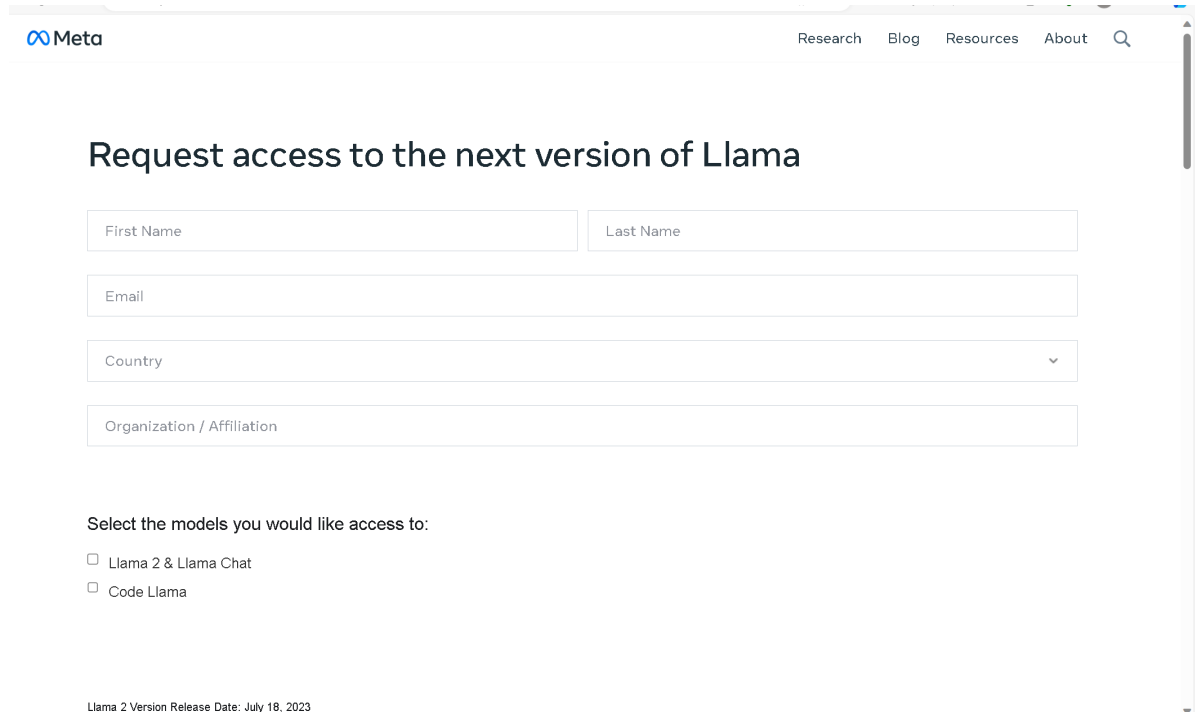
2. 下载llama-2

llama2可以在Meta官网和Huggingface网站上获取，为了安装的便捷性。我们最好遵从以下 workflow 完成：

1. Meta官网，使用美国节点+Gmail注册，并申请使用llama2的权限。
2. Huggingface官网，使用1中的gmail注册，并在Huggingface界面再次申请llama2权限。
3. 两个权限都申请通过以后，就可以开始准备下载环节。

注意：这里的两次注册最好使用同一个邮箱，且确保自己在美国节点的代理条件之下，预计两次审核时间不会超过一个小时。

Meta官网在完成如下表格的申请后，将会发送反馈邮件：

The image shows a screenshot of the Meta website's 'Request access to the next version of Llama' form. The form is located on the right side of the page, under the 'Research' tab. It contains several input fields: 'First Name', 'Last Name', 'Email', 'Country' (a dropdown menu), and 'Organization / Affiliation'. Below these fields, there is a section titled 'Select the models you would like access to:' with two checkboxes: 'Llama 2 & Llama Chat' and 'Code Llama'. At the bottom of the form, there is a small text indicating 'Llama 2 Version Release Date: July 18, 2023'.

huggingface在申请完成后，将会出现以下界面：

The screenshot shows the Hugging Face interface for the model **meta-llama/Llama-2-7b-chat-hf**. The main content area displays the model card, which includes a description of Llama 2 as a collection of pretrained and fine-tuned generative text models. The sidebar on the left contains navigation links like 'Model card', 'Files and versions', and 'Llama 2'. The right-hand panel provides additional details, including the number of downloads (950,482), the model size (6.74B params), and a list of spaces using the model.

在申请完成之后，我们就可以下载Huggingface的llama2版本。

Read [our paper](#), learn more [about the model](#), or get started with [code on GitHub](#).

[Llama Model Index](#)

Model	Llama2	Llama2-hf	Llama2-chat	Llama2-chat-hf
7B	Link	Link	Link	Link
13B	Link	Link	Link	Link
70B	Link	Link	Link	Link

llama2有许多版本，如上图，7B 13B是参数规模，不确定计算资源能力的情况下，先选择7B版本，实测3卡4090可以轻松运行7B 版本。

后缀主要是功能和版本的不同，如本文档选择的llama2-7b-chat-hf便是7B参数规模的对话机器人，自带使用Gradio构建的简易交互界面。

考虑到网络问题，我们需要使用的是Huggingface的镜像站和Hugg官方下载工具 `huggingface-cli`

镜像站在：hf-mirror.com - [Huggingface 镜像站](#)

按照以下方式执行，可以得到下载在当前目录下的具体文件。

```

1 pip install -U huggingface_hub # 安装
2
3 export HF_ENDPOINT=https://hf-mirror.com #切换镜像源，每次使用huggingface-cli镜像
  源下载都需要提前执行这一条命令
4
5 huggingface-cli download --resume-download meta-llama/Llama-2-7b-chat-hf --
  local-dir Llama-2-7b-chat-hf --local-dir-use-symlinks False --token
  hf_**Your token** # 替换自己的Huggingface token后，即可下载到本地，注意参数--local-
  dir-use-symlinks 必须添加，不然大于5M的文件会被自动生成软链接，后续llama2无法启动
6
7

```

这样我们就下载完成了llama2-7b-chat-hf。

3. Docker运行

本文档运行环境为docker,优点是环境隔离足够强，且由于docker对端口需要映射，可以解决后续一个很麻烦的端口问题。

(目前开源模型，llama2和stable diffusion等等，开源仓库都使用7860端口作为交互端口，如果服务器7860端口被占用，就需要腾出端口或者修改工程文件来实现。但是修改工程文件需要花费更多时间寻找对应工程文件，docker的端口映射可以帮助我们很简单地解决这个问题。)

要在docker环境下运行，自然需要安装docker,具体可以参考参考链接一节。自己配置docker file和环境依赖很耗时间，本着不重复造轮子的原则，使用前辈已经配置好的docker 库文件来完成。

库文件地址：[soulteary/docker-llama2-chat: Play LLaMA2 \(official / 中文版 / INT4 / llama2.cpp\). Together! ONLY 3 STEPS! \(non GPU / 5GB vRAM / 8~14GB vRAM\) \(github.com\)](https://github.com/soulteary/docker-llama2-chat: Play LLaMA2 (official / 中文版 / INT4 / llama2.cpp). Together! ONLY 3 STEPS! (non GPU / 5GB vRAM / 8~14GB vRAM) (github.com))

该库已经有安装说明，可在参考链接中查阅。由于本文档定位为从0到1的问题，所以对该库文件的文件不作修改，力求最快速度跑通。

下载github 仓库，需要使用镜像站，可以使用kkgithub.com来下载。由于镜像站稳定性差，建议搜索需要下载时可以使用的镜像站，并参考对应命令来下载。也可以下载到本地后，使用vscode 配置SFTP上传文件。

下面我们需要完成一系列调整工作：

1. 调整目录结构

```

1 # 创建一个新的目录，用于存放我们的模型
2 mkdir meta-llama
3
4 # 将下载好的模型移动到目录中
5 mv Llama-2-7b-chat-hf meta-llama/
6 mv Llama-2-13b-chat-hf meta-llama/
7 mv Llama-2-70b-chat-hf meta-llama/
8
9

```

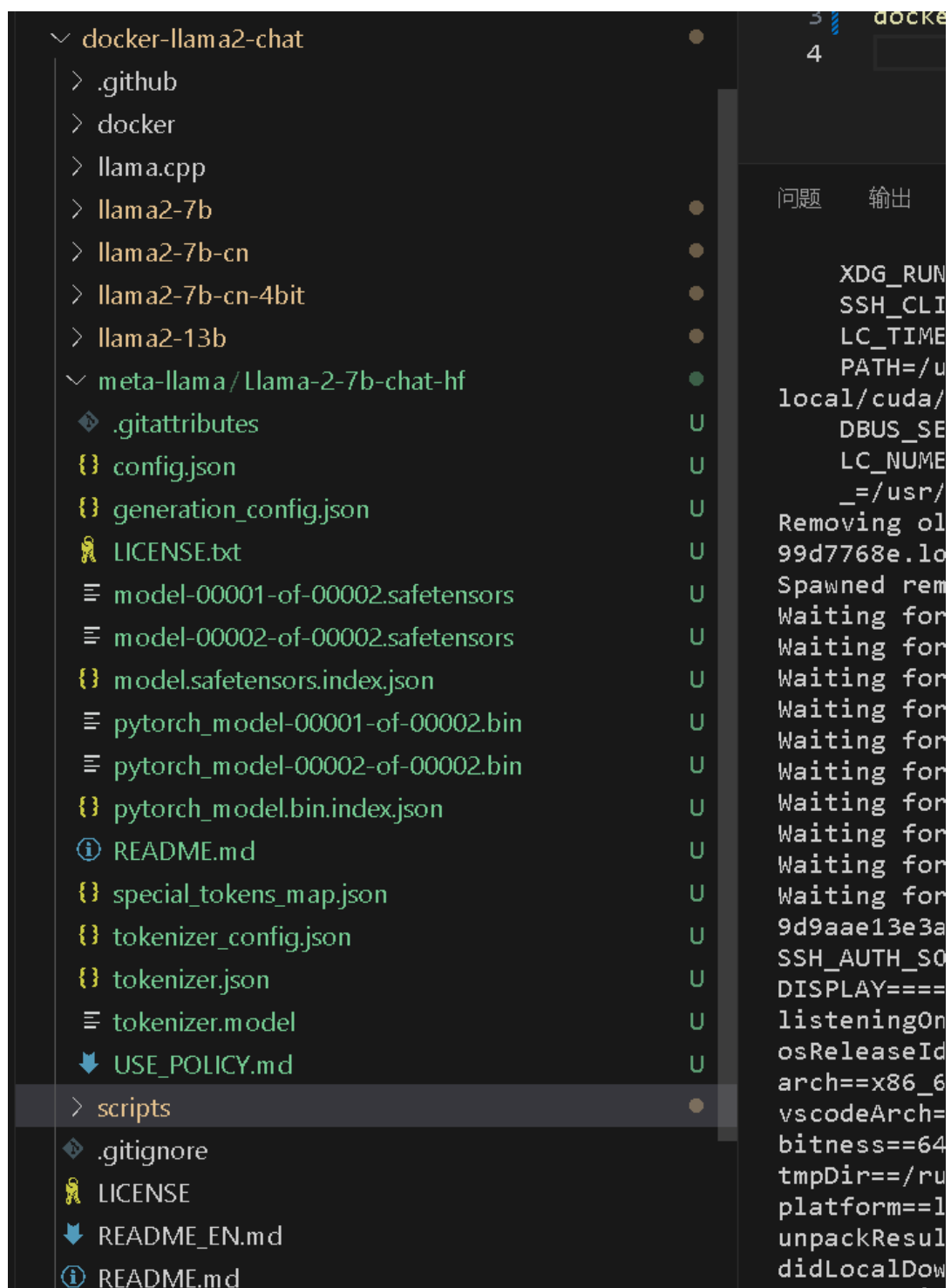
2. 完整的目录结构类似下面这样，所有的模型都在我们创建的 meta-llama 目录的下一级中

```

1 # tree -L 2 meta-llama
2 meta-llama
3 |— Llama-2-13b-chat-hf
4 |   |— added_tokens.json

```

```
5 | | └─ config.json
6 | | └─ generation_config.json
7 | | └─ LICENSE.txt
8 | | └─ model-00001-of-00003.safetensors
9 | | └─ model-00002-of-00003.safetensors
10 | | └─ model-00003-of-00003.safetensors
11 | | └─ model.safetensors.index.json
12 | | └─ pytorch_model-00001-of-00003.bin
13 | | └─ pytorch_model-00002-of-00003.bin
14 | | └─ pytorch_model-00003-of-00003.bin
15 | | └─ pytorch_model.bin.index.json
16 | | └─ README.md
17 | | └─ Responsible-Use-Guide.pdf
18 | | └─ special_tokens_map.json
19 | | └─ tokenizer_config.json
20 | | └─ tokenizer.model
21 | | └─ USE_POLICY.md
22 | └─ Llama-2-7b-chat-hf
23 |   └─ added_tokens.json
24 |   └─ config.json
25 |   └─ generation_config.json
26 |   └─ LICENSE.txt
27 |   └─ model-00001-of-00002.safetensors
28 |   └─ model-00002-of-00002.safetensors
29 |   └─ model.safetensors.index.json
30 |   └─ models--meta-llama--Llama-2-7b-chat-hf
31 |   └─ pytorch_model-00001-of-00003.bin
32 |   └─ pytorch_model-00002-of-00003.bin
33 |   └─ pytorch_model-00003-of-00003.bin
34 |   └─ pytorch_model.bin.index.json
35 |   └─ README.md
36 |   └─ special_tokens_map.json
37 |   └─ tokenizer_config.json
38 |   └─ tokenizer.json
39 |   └─ tokenizer.model
40 |   └─ USE_POLICY.md
41
```



开始构建并运行：

```

1  # 进入程序目录
2  cd docker-llama2-chat
3  # 构建 7B 镜像
4  bash scripts/make-7b.sh
5  # 或者，构建 13B 镜像
6  bash scripts/make-13b.sh
7
8  #构建完成后
9
10 # 运行 7B 镜像，应用程序
11 bash scripts/run-7b.sh
12 # 或者，运行 13B 镜像，应用程序
13 bash scripts/run-13b.sh
14

```

```

15 #命令执行后，如果一切顺利，你将看到类似下面的日志：
16
17 =====
18 == PyTorch ==
19 =====
20
21 NVIDIA Release 23.06 (build 63009835)
22 PyTorch Version 2.1.0a0+4136153
23
24 Container image Copyright (c) 2023, NVIDIA CORPORATION & AFFILIATES. All
    rights reserved.
25
26 Copyright (c) 2014-2023 Facebook Inc.
27 Copyright (c) 2011-2014 Idiap Research Institute (Ronan Collobert)
28 Copyright (c) 2012-2014 Deepmind Technologies (Koray Kavukcuoglu)
29 Copyright (c) 2011-2012 NEC Laboratories America (Koray Kavukcuoglu)
30 Copyright (c) 2011-2013 NYU (Clement Farabet)
31 Copyright (c) 2006-2010 NEC Laboratories America (Ronan Collobert, Leon
    Bottou, Iain Melvin, Jason Weston)
32 Copyright (c) 2006 Idiap Research Institute (Samy Bengio)
33 Copyright (c) 2001-2004 Idiap Research Institute (Ronan Collobert, Samy
    Bengio, Johnny Mariethoz)
34 Copyright (c) 2015 Google Inc.
35 Copyright (c) 2015 Yangqing Jia
36 Copyright (c) 2013-2016 The Caffe contributors
37 All rights reserved.
38
39 Various files include modifications (c) NVIDIA CORPORATION & AFFILIATES. All
    rights reserved.
40
41 This container image and its contents are governed by the NVIDIA Deep
    Learning Container License.
42 By pulling and using the container, you accept the terms and conditions of
    this license:
43 https://developer.nvidia.com/ngc/nvidia-deep-learning-container-license
44
45 WARNING: CUDA Minor Version Compatibility mode ENABLED.
46 Using driver version 525.105.17 which has support for CUDA 12.0. This
    container
47 was built with CUDA 12.1 and will be run in Minor Version Compatibility
    mode.
48 CUDA Forward Compatibility is preferred over Minor Version Compatibility
    for use
49 with this container but was unavailable:
50 [[Forward compatibility was attempted on non supported HW
    (CUDA_ERROR_COMPAT_NOT_SUPPORTED_ON_DEVICE) cuInit()=804]]
51 See https://docs.nvidia.com/deploy/cuda-compatibility/ for details.
52
53 Loading checkpoint shards: 100%|
    ████████████████████████████████████████████████████████████████████████████
    ████████████████████████████████████████████████████████████████████████████ | 2/2
    [00:05<00:00, 2.52s/it]
54 Caching examples at: '/app/gradio_cached_examples/20'
55 Caching example 1/5

```

```

56 /usr/local/lib/python3.10/dist-
packages/transformers/generation/utils.py:1270: UserWarning: You have
modified the pretrained model configuration to control generation. This is a
deprecated strategy to control generation and will be removed soon, in a
future version. Please use a generation configuration file (see
https://huggingface.co/docs/transformers/main_classes/text_generation )
57     warnings.warn(
58 Caching example 2/5
59 Caching example 3/5
60 Caching example 4/5
61 Caching example 5/5
62 Caching complete
63
64 /usr/local/lib/python3.10/dist-packages/gradio/utils.py:839: UserWarning:
Expected 7 arguments for function <function generate at 0x7f3e096a1000>,
received 6.
65     warnings.warn(
66 /usr/local/lib/python3.10/dist-packages/gradio/utils.py:843: UserWarning:
Expected at least 7 arguments for function <function generate at
0x7f3e096a1000>, received 6.
67     warnings.warn(
68 Running on local URL:  http://0.0.0.0:7860
69
70 To create a public link, set `share=True` in `launch()`.
71
72

```

然后使用浏览器打开<https://localhost:7860>即可运行。

如果是使用vscode远程开发，会先出现下图的提示：

```

docker-llama2-chat > scripts > $ run-7b.sh
1
2
3 k=-1 --ulimit stack=67108864 --rm -it -v `pwd`/meta-llama:/app/meta-llama -p 7865:7860 soulteary/llama2:7b
4

```

```

• (base) fuhongye@amax:~$ ls
docker-llama2-chat llama2 llama2-2-7b-hf
• (base) fuhongye@amax:~$ cd docker*
○ (base) fuhongye@amax:~/docker-llama2-chat$ bash scripts/run-7b.sh

=====
== PyTorch ==
=====

NVIDIA Release 23.06 (build 63009835)
PyTorch Version 2.1.0a0+4136153

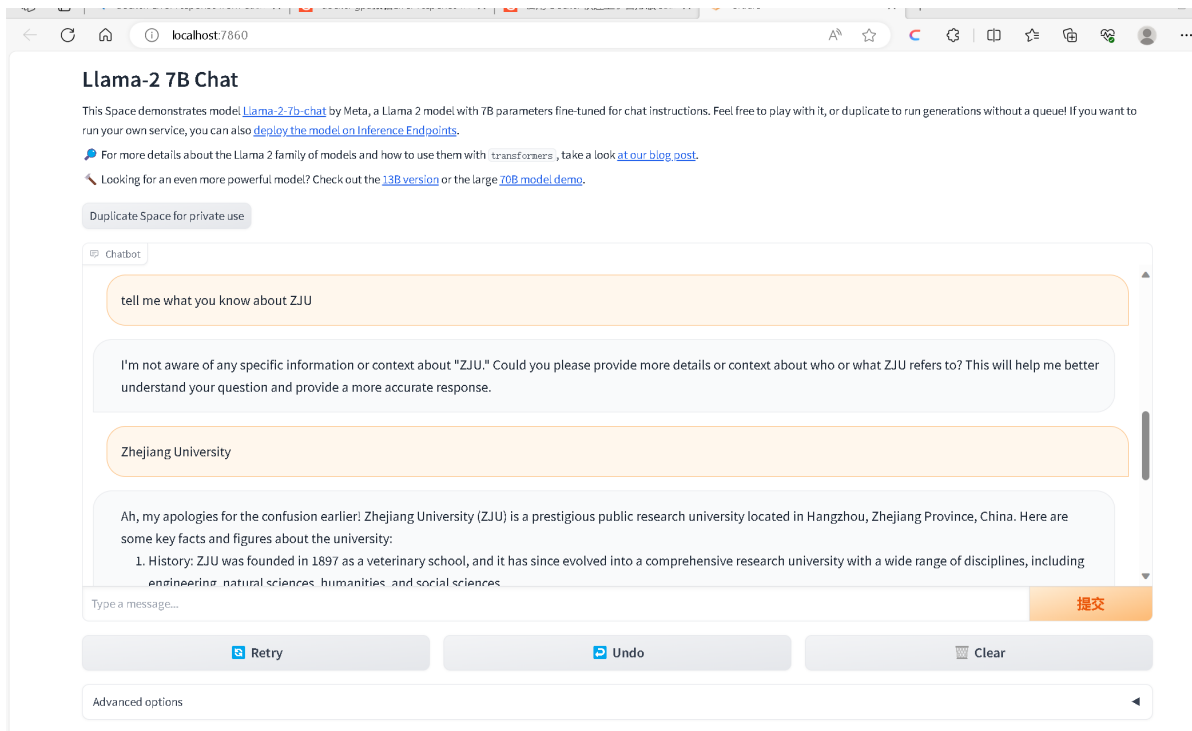
Container image Copyright (c) 2023, NVIDIA CORPORATION & AFFILIATES. All rights reserved.

Copyright (c) 2014-2023 Facebook Inc.
Copyright (c) 2011-2014 Idiap Research Institute (Ronan Collobert)
Copyright (c) 2012-2014 Deepmind Technologies (Koray Kavukcuoglu)
Copyright (c) 2011-2012 NEC Laboratories America (Koray Kavukcuoglu)

```

如果端口7860不可用，则可以修改为7865或者其他端口，如上图就修改为外部7865端口映射docker运行的7860端口。

然后vscode会提示是否需要转发端口上的内容，可以添加7865端口进入转发设置，然后打开任意浏览器使用llama2。



4. 参考链接

1. [最详细的ubuntu 安装 docker教程 - 知乎 \(zhihu.com\)](#)
2. [使用 Docker 快速上手官方版 LLaMA2 开源大模型_soulteary的博客-CSDN博客](#)
3. [【Nvidia-smi报错】Failed to initialize NVML: Driver/library version mismatch-CSDN博客](#)
4. [\[docker gpu报错Error response from daemon: could not select device driver "" with capabilities: \[\[gpu\]\]_Adenialzz的博客-CSDN博客\]\(https://blog.csdn.net/weixin_44966641/article/details/123760614\)](#)
5. [hf-mirror.com - Huggingface 镜像站](#)

written by Hongye Fu 2023.11.24, 感谢soulteary前辈的文档和仓库