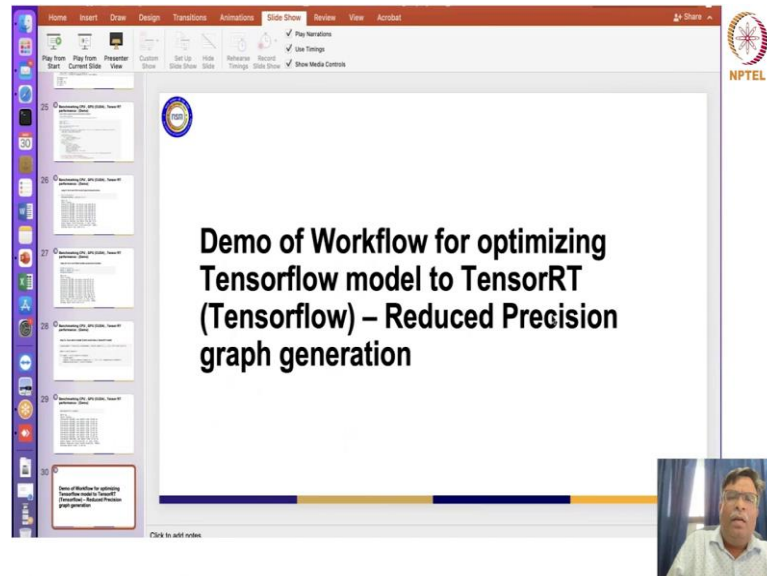


Applied Accelerated Artificial Intelligence
Prof. Satyadhyan Chickerur
Department of Computer Science and Engineering
Indian Institute of Technology, Palakkad

Lecture - 43

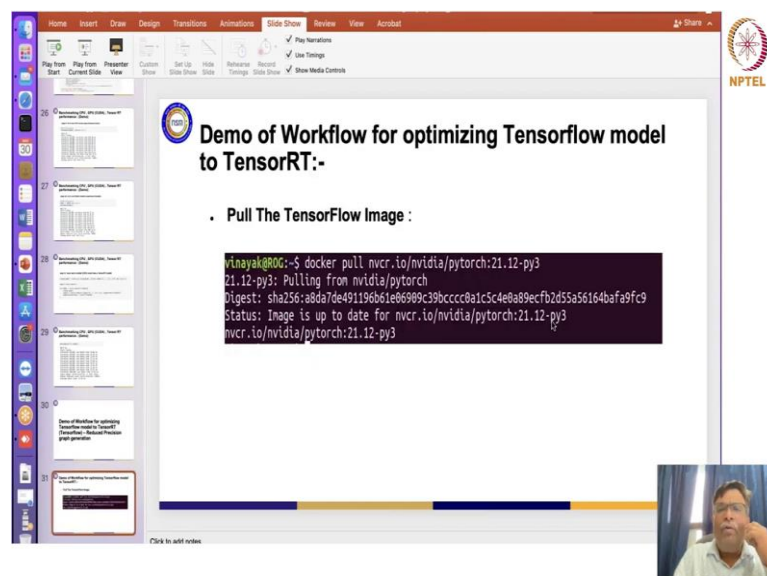
Accelerating neural network inference in PyTorch and TensorFlow Part 2

(Refer Slide Time: 00:15)

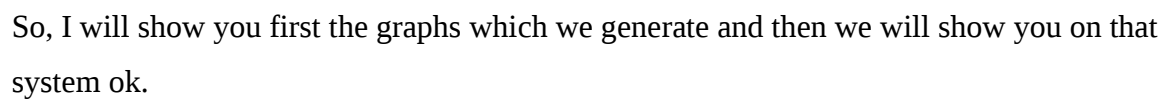


Now, one of the things which we would like to show you is a demo of workflow for optimizing the tensorflow model to TensorRT model using reduced precision right.

(Refer Slide Time: 00:29)



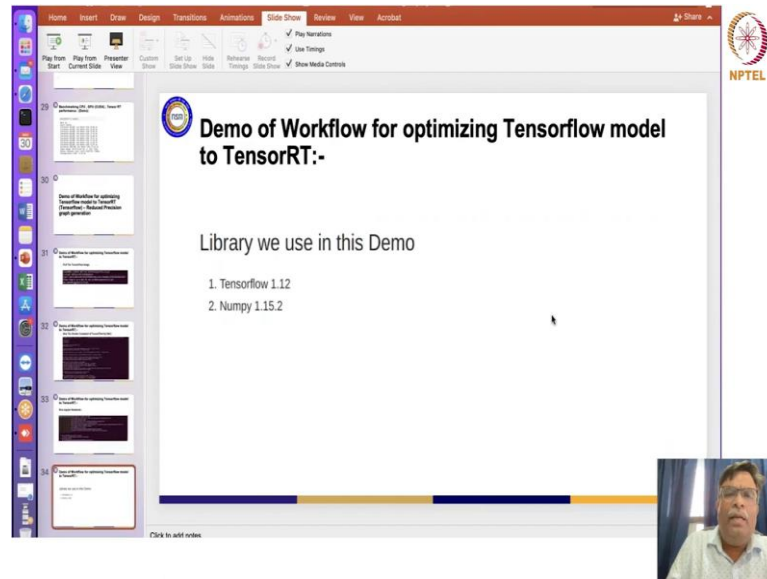
(Refer Slide Time: 00:38)



(Refer Slide Time: 00:46)



(Refer Slide Time: 00:48)

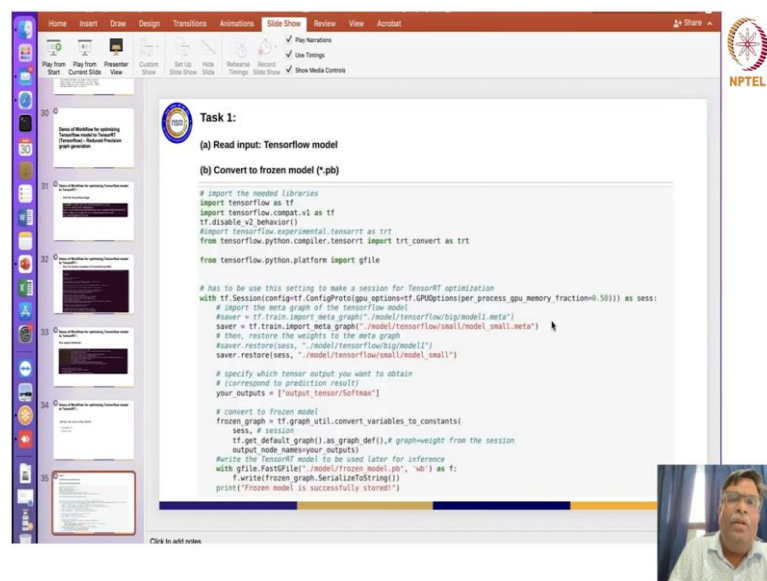


Demo of Workflow for optimizing Tensorflow model to TensorRT:-

Library we use in this Demo

1. Tensorflow 1.12
2. Numpy 1.15.2

(Refer Slide Time: 00:52)



Task 1:

(a) Read input: Tensorflow model
(b) Convert to frozen model (*.pb)

```
# Import the needed libraries
import tensorflow as tf
import tensorflow.compat.v1 as tf
tf.disable_v2_behavior()
import tensorflow.experimental.tensorrt as trt
from tensorflow.python.compiler.tensorrt import trt_convert as trt
from tensorflow.python.platform import gfile

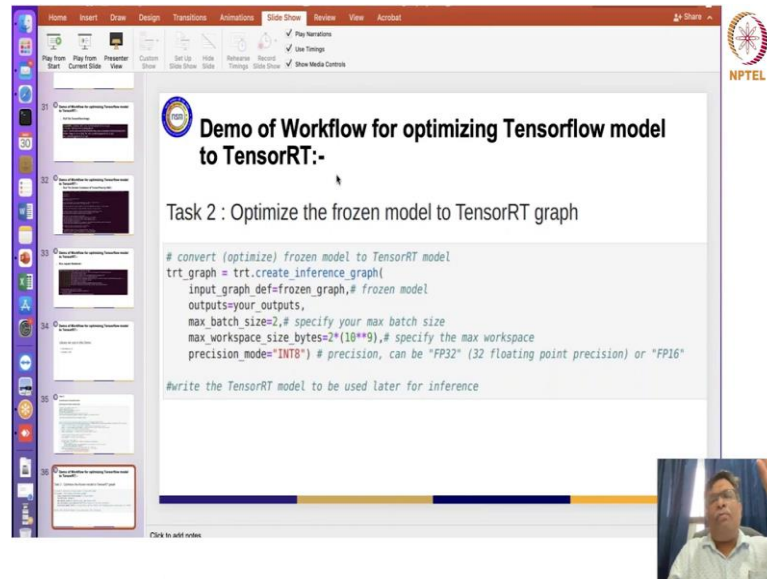
# has to be use this setting to make a session for TensorRT optimization
with tf.Session(config=tf.ConfigProto(gpu_options=tf.GPUOptions(per_process_gpu_memory_fraction=0.5))) as sess:
    # Import the meta graph of the tensorflow model
    saver = tf.train.import_meta_graph("./model/tensorflow/ship/model1.meta")
    # then, restore the weights to the meta graph
    saver.restore(sess, "./model/tensorflow/ship/model1")
    saver.restore(sess, "./model/tensorflow/ship/model1")

    # specify which tensor output you want to obtain
    # (correspond to prediction result)
    your_outputs = ["output_tensor_softmax"]

    # convert to frozen model
    frozen_graph = tf.graph_util.convert_variables_to_constants(
        sess, # session
        tf.get_default_graph().as_graph_def(), # graph=weight from the session
        output_names=your_outputs)
    # Write the TensorRT model to be used later for inference
    with gfile.FastFile("./model/frozen_model.pb", "wb") as f:
        f.write(frozen_graph.SerializeToString())
    print("Frozen model is successfully stored")
```

So, this is basically trying to pull the containers. Then you basically use this and then once you read a tensorflow model you convert to a frozen model which is *.pb ok this is what basically is a frozen model.

(Refer Slide Time: 01:07)



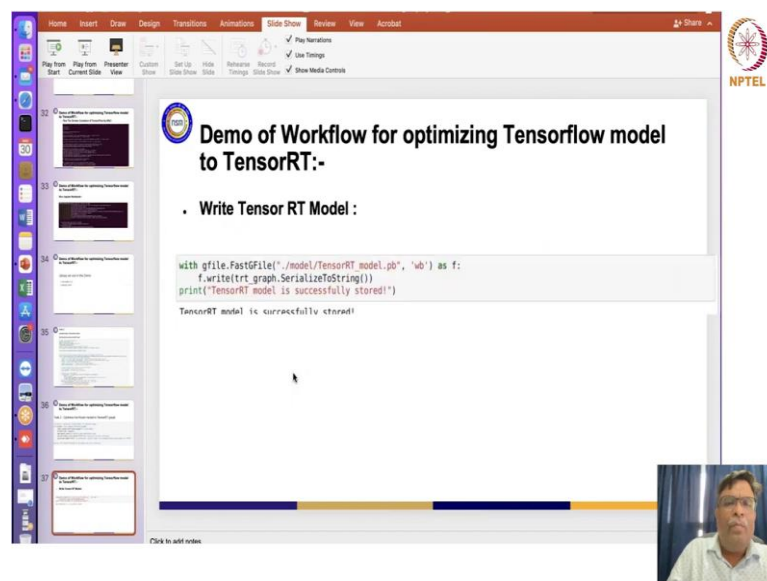
The screenshot shows a presentation slide with the title "Demo of Workflow for optimizing Tensorflow model to TensorRT:-". Below the title, it says "Task 2 : Optimize the frozen model to TensorRT graph". The slide contains a Python code block for converting a frozen TensorFlow model to a TensorRT model. The code is as follows:

```
# convert (optimize) frozen model to TensorRT model
trt_graph = trt.create_inference_graph(
    input_graph_def=frozen_graph, # frozen model
    outputs=your_outputs,
    max_batch_size=2, # specify your max batch size
    max_workspace_size_bytes=2*(10**9), # specify the max workspace
    precision_mode="INT8" # precision, can be "FP32" (32 floating point precision) or "FP16"

#write the TensorRT model to be used later for inference
```

A small video inset in the bottom right corner shows a man speaking.

(Refer Slide Time: 01:27)



The screenshot shows a presentation slide with the title "Demo of Workflow for optimizing Tensorflow model to TensorRT:-". Below the title, it says "Write Tensor RT Model :". The slide contains a Python code block for saving the TensorRT model. The code is as follows:

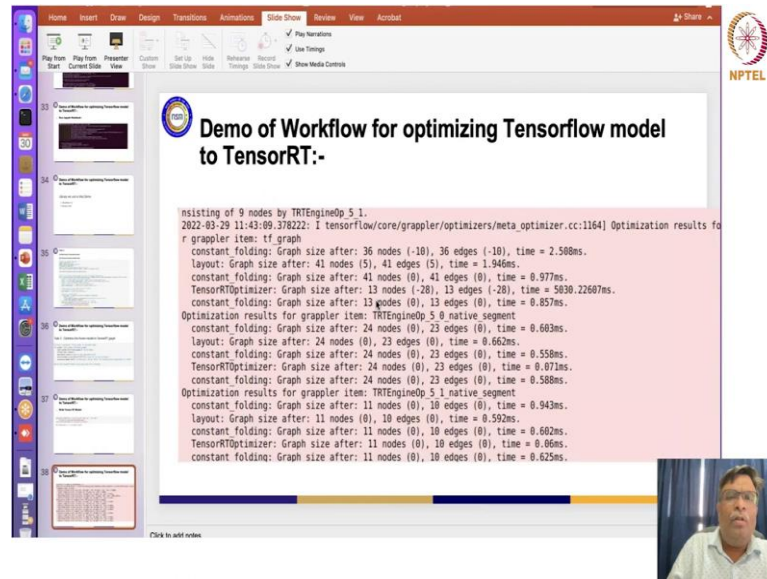
```
with gfile.GFile("./model/tensorrt_model.pb", "wb") as f:
    f.write(trt_graph.SerializeToString())
    print("TensorRT model is successfully stored!")

TensorRT model is successfully stored!
```

A small video inset in the bottom right corner shows a man speaking.

And then you optimize that frozen model to a TensorRT graph right this is how you do it you convert the frozen model to a TensorRT graph using trt graph. Then you actually put trt create inference graph all of these parameters right. And the precision mode which we discussed ok and everything like that. And then you basically get this something like a TensorRT model this successfully stored because we are printing it because we are trying to store this model ok.

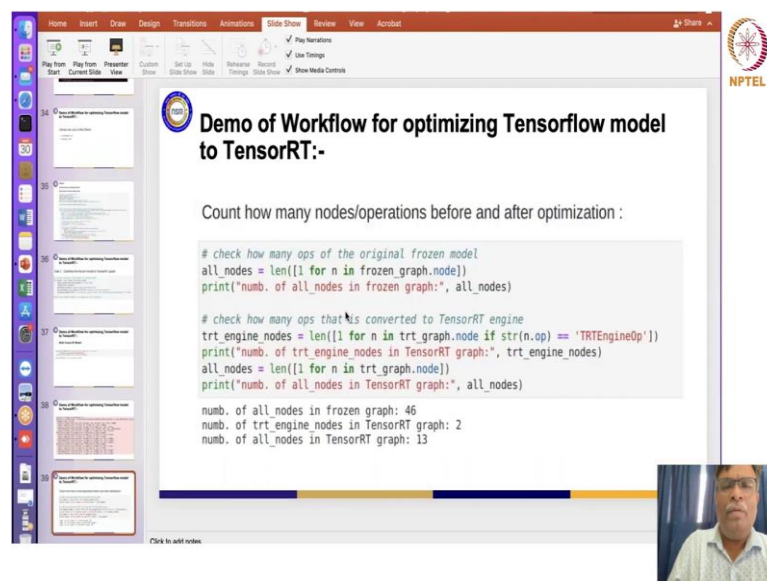
(Refer Slide Time: 01:38)



Demo of Workflow for optimizing Tensorflow model to TensorRT:-

nsisting of 9 nodes by TRTEngineOp 5.1.
2022-03-29 11:43:09.378222: I tensorflow/core/grappler/optimizers/meta_optimizer.cc:1164] Optimization results for
r grappler item: tf graph
constant folding: Graph size after: 36 nodes (-10), 36 edges (-10), time = 2.580ms.
layout: Graph size after: 41 nodes (5), 41 edges (5), time = 1.946ms.
constant folding: Graph size after: 41 nodes (0), 41 edges (0), time = 0.977ms.
TensorRTOptimizer: Graph size after: 13 nodes (-28), 13 edges (-28), time = 5830.22607ms.
constant folding: Graph size after: 13 nodes (0), 13 edges (0), time = 0.857ms.
Optimization results for grappler item: TRTEngineOp 5.0 native segment
constant folding: Graph size after: 24 nodes (0), 23 edges (0), time = 0.603ms.
layout: Graph size after: 24 nodes (0), 23 edges (0), time = 0.662ms.
constant folding: Graph size after: 24 nodes (0), 23 edges (0), time = 0.558ms.
TensorRTOptimizer: Graph size after: 24 nodes (0), 23 edges (0), time = 0.071ms.
constant folding: Graph size after: 24 nodes (0), 23 edges (0), time = 0.588ms.
Optimization results for grappler item: TRTEngineOp 5.1 native segment
constant folding: Graph size after: 11 nodes (0), 10 edges (0), time = 0.943ms.
layout: Graph size after: 11 nodes (0), 10 edges (0), time = 0.592ms.
constant folding: Graph size after: 11 nodes (0), 10 edges (0), time = 0.602ms.
TensorRTOptimizer: Graph size after: 11 nodes (0), 10 edges (0), time = 0.06ms.
constant folding: Graph size after: 11 nodes (0), 10 edges (0), time = 0.625ms.

(Refer Slide Time: 01:55)



Demo of Workflow for optimizing Tensorflow model to TensorRT:-

Count how many nodes/operations before and after optimization :

```
# check how many ops of the original frozen model
all_nodes = len([l for n in frozen_graph.node])
print("numb. of all_nodes in frozen graph:", all_nodes)

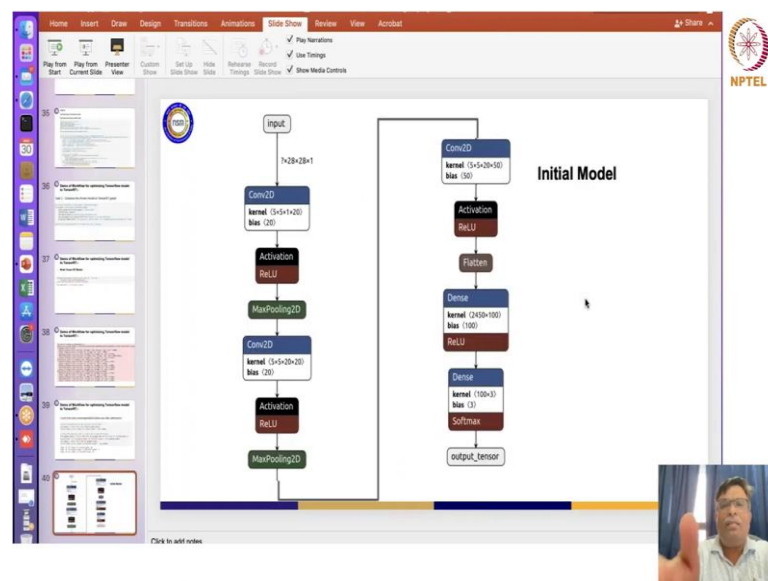
# check how many ops that is converted to TensorRT engine
trt_engine_nodes = len([l for n in trt_graph.node if str(n.op) == 'TRTEngineOp'])
print("numb. of trt_engine_nodes in TensorRT graph:", trt_engine_nodes)
all_nodes = len([l for n in trt_graph.node])
print("numb. of all_nodes in TensorRT graph:", all_nodes)

numb. of all_nodes in frozen graph: 46
numb. of trt_engine_nodes in TensorRT graph: 2
numb. of all_nodes in TensorRT graph: 13
```

So, this is how we actually when you run you get something like the tf graph ok. What is the graph size? How many nodes? Ok. What is the graph size everything like this we will show you gradually of how actually it is optimized and ultimately what you get in the end is something like this. Number of all nodes in the frozen graph is 46 number of trt engines nodes in the TensorRT graph is 2 we will see this ok and number of all nodes in the TensorRT graph is 13.

So, this basically means you are trying to correlate a model which you have actually got it from a tensorflow ok to something which is optimized like this ok. So, how many nodes in the frozen graph? 46. Now what you do? You convert that into a TensorRT graph with how many nodes? 13 ok. And how many engine nodes in the TensorRT graph? There are only 2. We will see what does it mean by engine node and what are various graphs. So, why there are various nodes in the graph, but this is how it is.

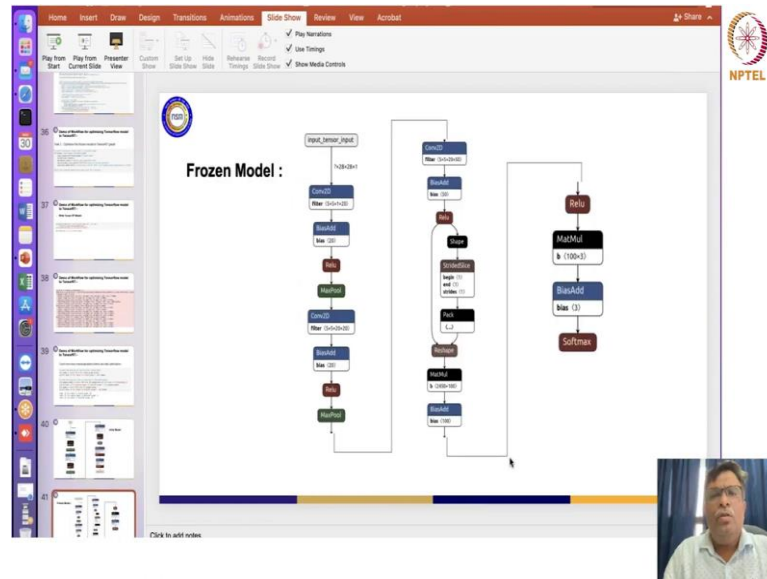
(Refer Slide Time: 02:46)



So, let us try to understand this I hope at least on the tensor board you would have seen this tensorflow graphs right. So, this is how a graph for a tensorflow program looks like ok this is a classification program you have input you have convolution layer with this size, this size of the kernel.

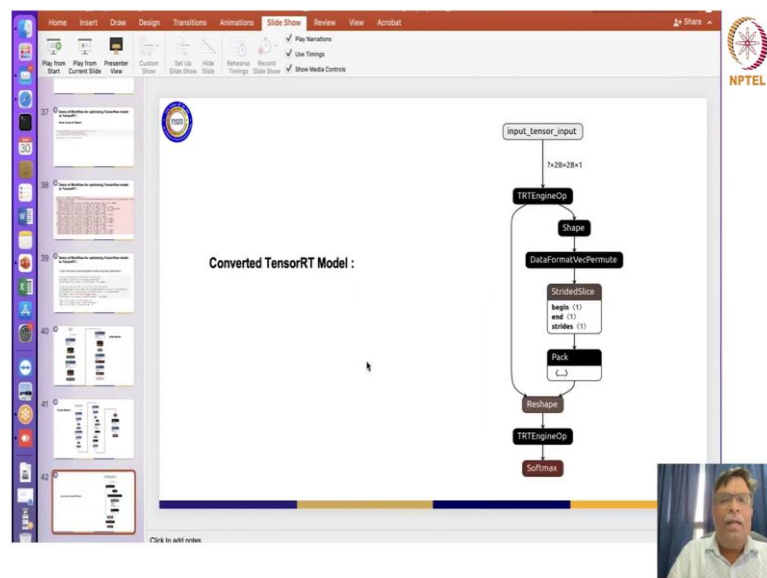
Then you have the ReLU activation then max pooling then again the convolution layer all of this like this is the initial model like you have a flattened thing you have soft max ok and then you have this output tensor. So, input and output tensor and this is the initial model ok. So, now, from this initial model after those steps you will get a frozen model. So, this is how frozen model is right.

(Refer Slide Time: 03:31)



So, you get a frozen model like this ok. This is converted ok this frozen model and then not going into the details, but as you go along you can actually understand right which of these layers ok are like basically you know finalized and frozen right. So, this and this is almost the same with only issue is you cannot change anything in this, this is not trainable now, but this is trained model this can be trained further right. And then the next thing is that you get a TensorRT model right.

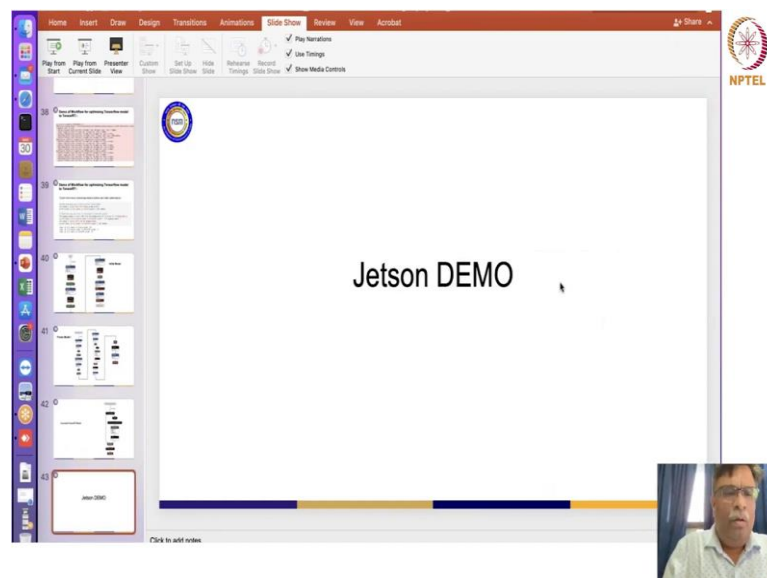
(Refer Slide Time: 04:14)



So, this is how the trt engine ok operator gets you this format. So, this particular model ok has been optimized to this wherein it uses fusing ok. Fusion of the layers horizontal fusion, vertical fusion, everything happens you get the tensor input ok. The same way and then you have this final output at the softmax layer this is what actually is converted from a frozen model to a TensorRT model.

We are not going into the details at present, but this is a very very good detailed what to say content ok which basically requires you a very very good understanding of how basically it has to be done if you understand the internal intricacies of the tensorflow the graph ok and all of that. So, we are just trying to show you in a short time as to how everything happens ok.

(Refer Slide Time: 05:30)



(Refer Slide Time: 05:43)

So let us try to actually do some demo now of how everything is done. So, let me just. So, basically we are going to import ok the tensorflow model this is how you actually import it ok. So, you are trying to import a model which is a small model ok. And then you basically try to actually convert that one minute convert that to a frozen model.

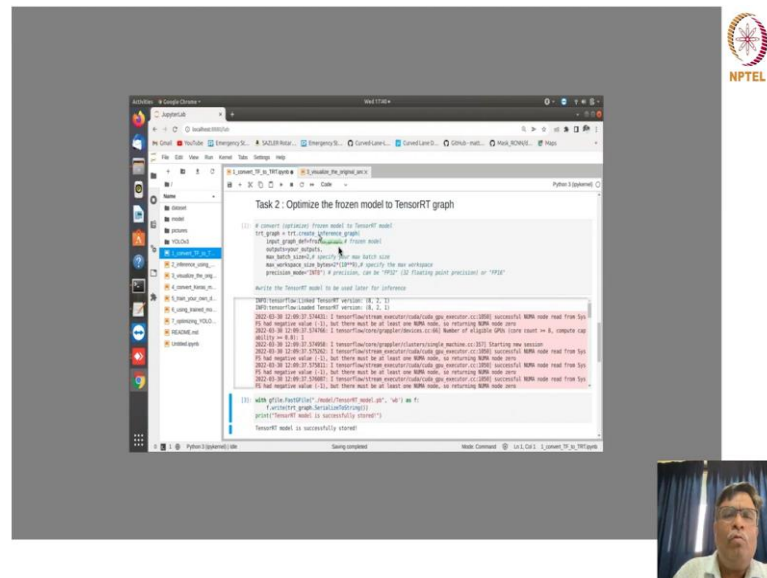
(Refer Slide Time: 06:14)

[illegible]

So, what you are trying to do is you are trying to actually generate a frozen graph ok from basically this. And if you see here you will actually try to use the what to say the frozen model ok is to be generated using CPU, GPU, CUDA, GPU executor right. So,

here what we are trying to do is we are trying to actually use a tensor flow. We are trying to understand which nvidia geforce gtx on which we are running ok on which device we are trying to convert it and then we get a frozen model here right.

(Refer Slide Time: 06:57)



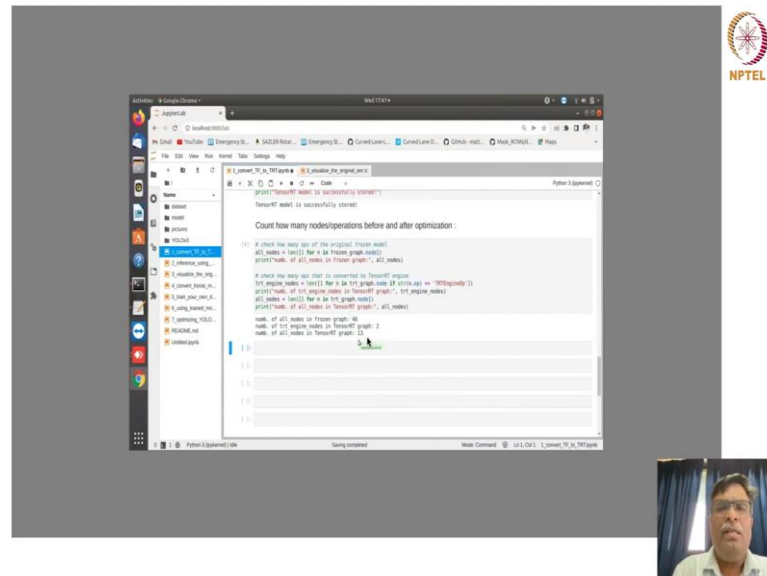
And when we try to optimize the frozen model to a TensorRT graph ok. So, this is a frozen model now. So, we are just trying to understand and generate a trt graph ok. So, we are going to create a inference graph and this is how the thing is my maximum batch size.

What batch size what is the precision we are trying to do ok. And then what is the input of to this particular graph definition. So, this is a frozen model right which we have used. So, we are using this frozen model in the previous step to generate a trt graph definition right and from that we are trying to create a inference graph right.



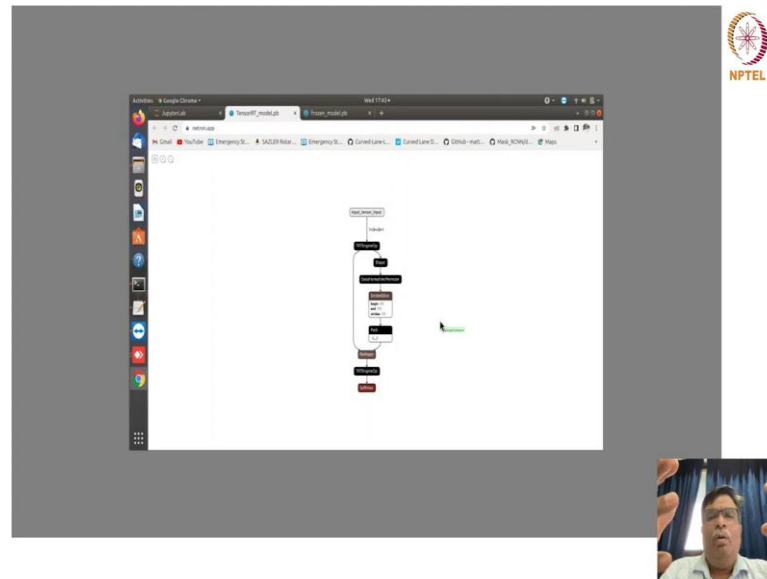
And then you basically try to actually ok understand right how is this graph size ok getting modified right. So, it is all linked to tensors and then the graphs right.

(Refer Slide Time: 08:47)



So, this is how it is going to actually do and then you basically ultimately ok. See a tensor model which is successively successfully stored and you can actually count ok how many nodes and how many operations were there before and after the operations this is what I showed you in the slide as well. So, number of nodes in the frozen graph is 46; number of nodes in the RT graph is 13 and then convert it into TensorRT model you can visualize this ok.

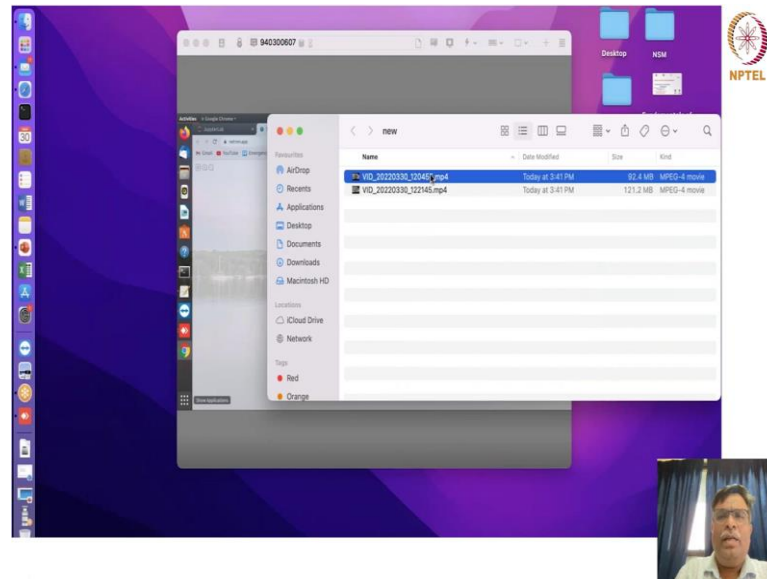
(Refer Slide Time: 09:22)



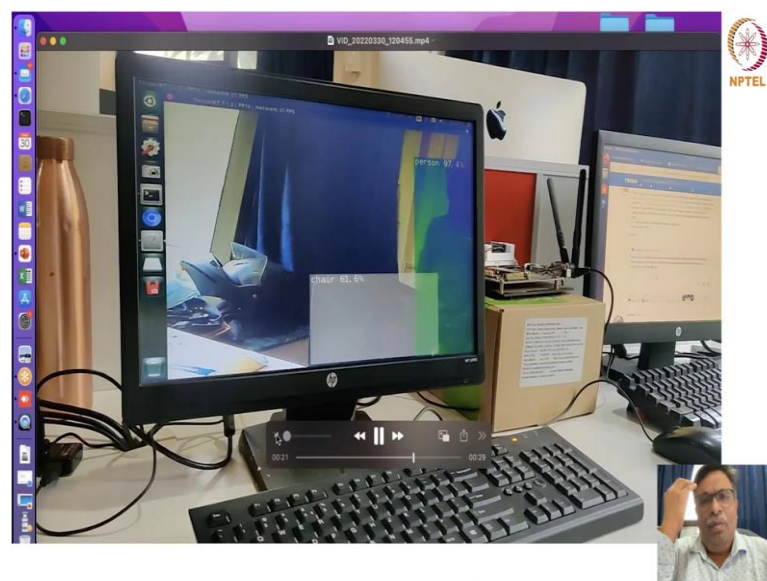
So, this is just to make you feel ok that once it is actually optimized for inferencing ok. How small the scale of the whole graph becomes ok and that is why the inferencing is such a serious issue wherein optimization is of utmost importance right. So, I hope this has made certain things clear on the part of how the optimization happens right. So, this is how the fusion happens and when we run the program.

Now we will show you right how the optimization or how the fusion happens right it shows you step by step ok. We will try to show you that now before that I will just try to run one or two videos in the meantime we will set that up for you to actually see ok. So, I will run the videos and in the meantime I will just set the camera also. So, give me a sec and I will try to run the video which we actually showed you yesterday certain applications.

(Refer Slide Time: 10:34)



(Refer Slide Time: 10:36)

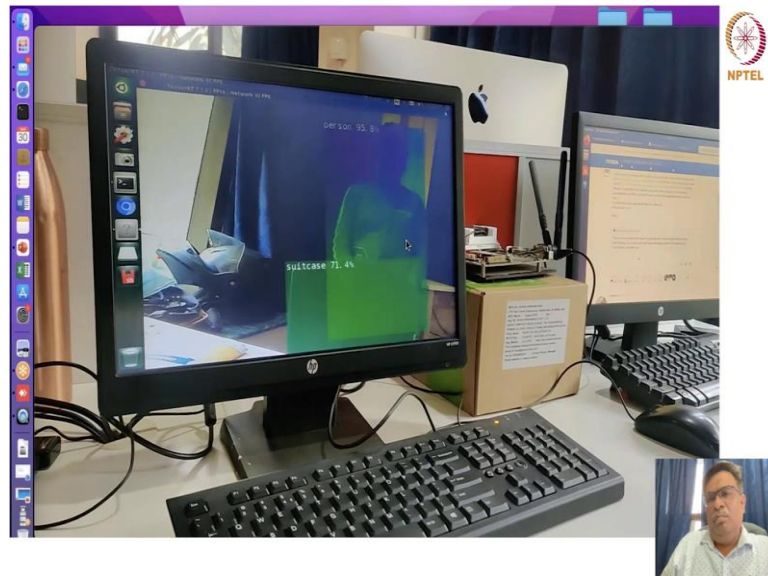


So, today we again try to do the same program run on a TensorRT optimized program ok on our jetson tx 1. So, you see this its a very very small program right it is not very what to say. Done for a specific application, but these are all available online we will give you the links. So, that you can also download and do.

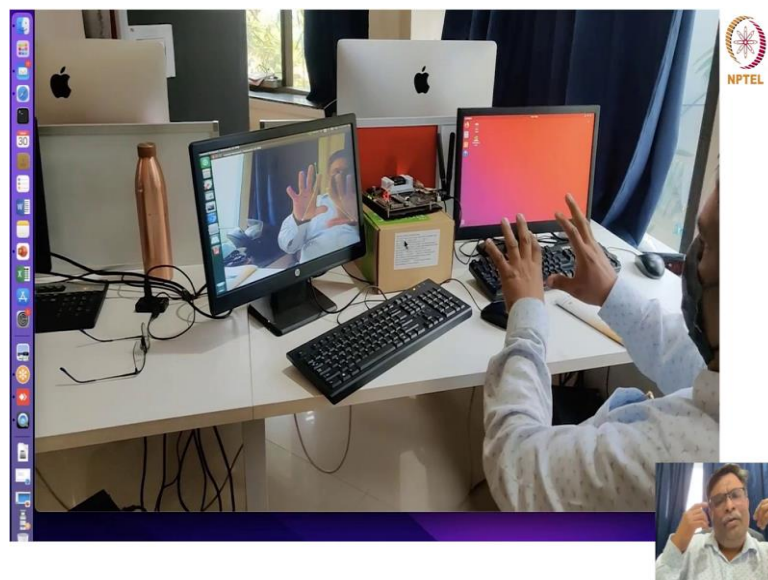
So, it can detect multiple persons multiple chairs ok not very much using a very very highly trained model, but ok. For some of the things right like suitcase and chair it is

showing, but ok these are certain things which you need to rectify, but we are just trying to show you that like this is how it can be done on a inferencing edge device right.

(Refer Slide Time: 11:38)



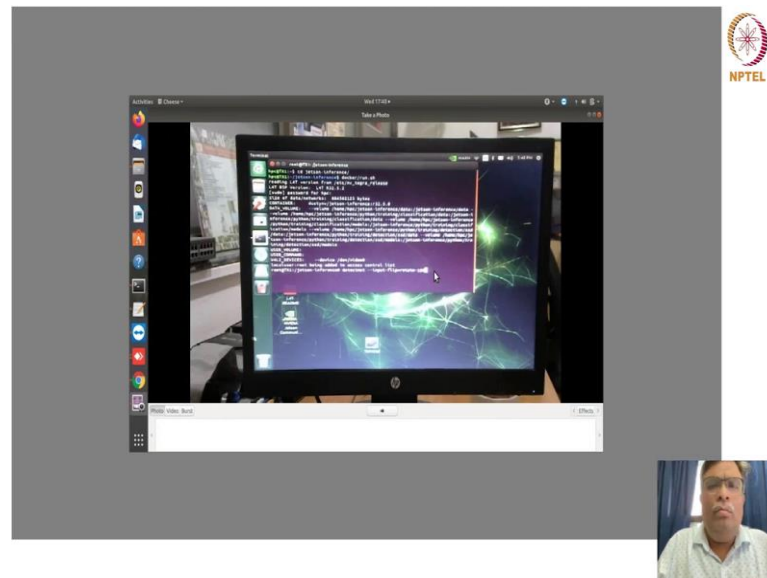
(Refer Slide Time: 11:49)



So, this is one thing which we try to show you and another one is wherein we do something of detection ok pose detection. So, this basically is a pose detection program which we are trying to run with a inbuilt onboard camera of on jetson tx 1. So, this is what we actually tried to do.

So, you have got a pose detection, you have got classification then you basically have a lot of other programs which are there ok. So, suppose got stuck or something, but yeah so something like this. So, let me just try to run it and 1 minute yes.

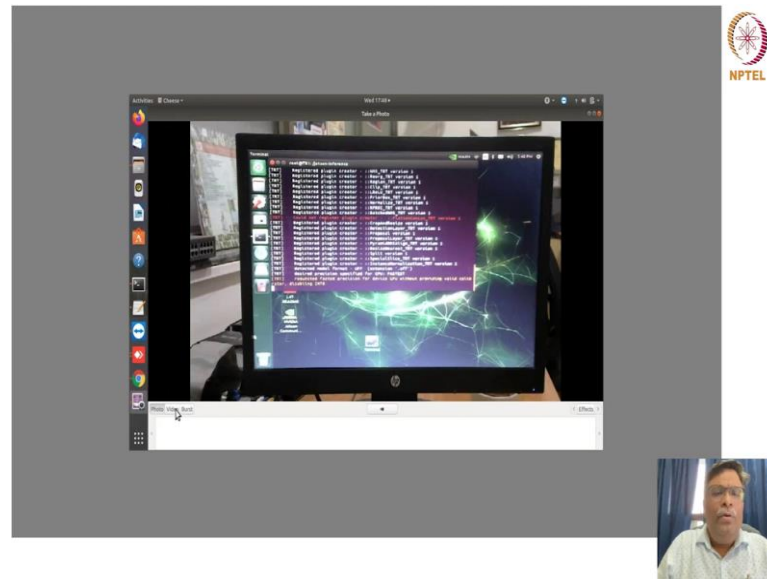
(Refer Slide Time: 12:54)



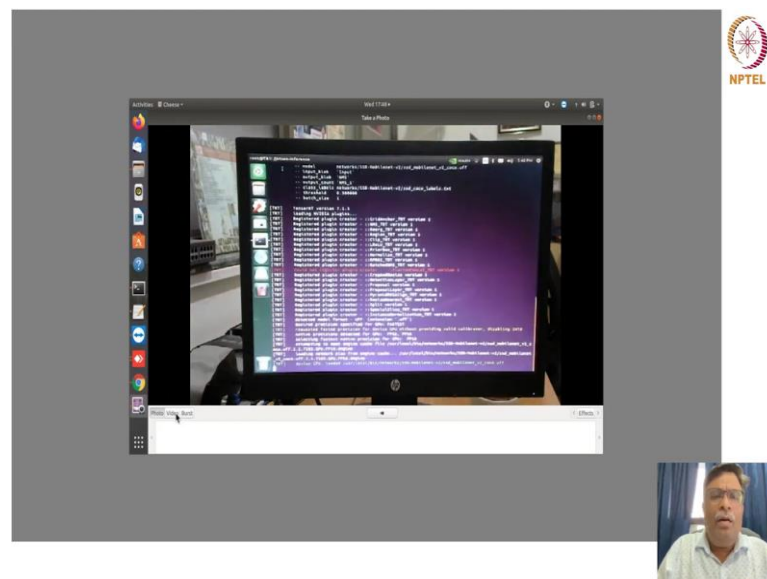
So, if you can see the monitor which basically is connected to the jetson tx 1. What we are trying to do is I do not know whether the resolution of the camera is good in a sense, but we have to enlarge it a bit. Let me check if we can enlarge it 1 minute no this is the maximum size let us see what happens it ok.

So, yeah somewhat it increase let me just check once again no ok. So, what we are trying to do is now we are trying to run detect net program ok for actually classifying right people ok and chair and all of that right. So, let us run the program and then see what basically happens ok image gets stuck it seems I do not know here we will have to actually convert it into video thing yeah.

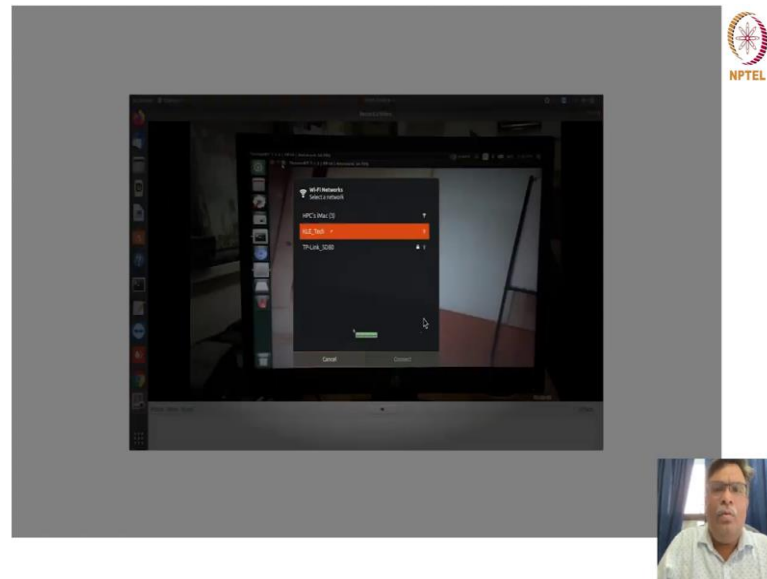
(Refer Slide Time: 14:20)



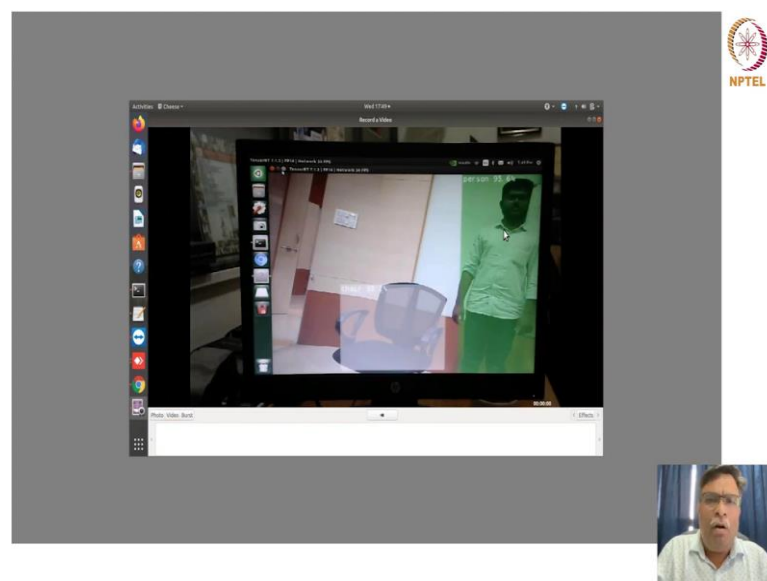
(Refer Slide Time: 14:31)



(Refer Slide Time: 14:34)



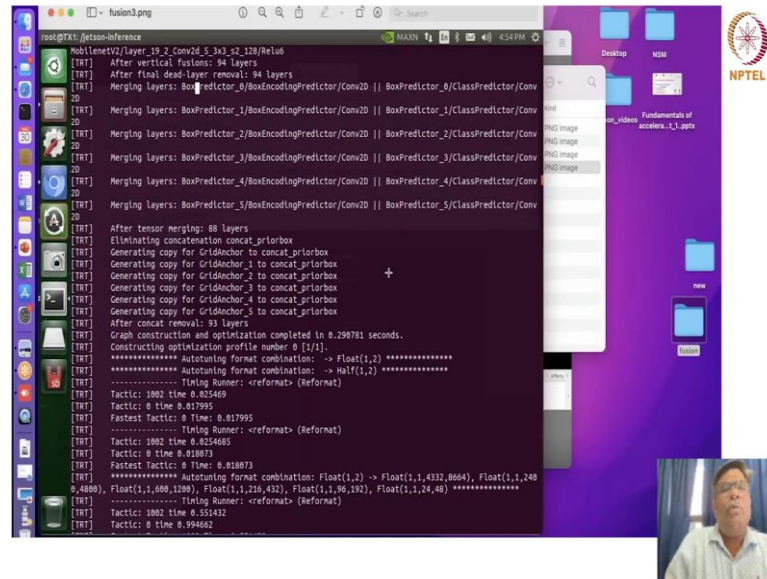
(Refer Slide Time: 14:40)



So, if you are seeing this you just convert ok. So, here if you are seeing this that this 1 minute it is a bit slow because of the lag, but the idea is that this is happening and then if we show you the conversion of things now the camera is not correct. So, it is not actually showing, but I suppose it is visible yeah.

So, if we stop this we can show you the scrolled history and there you can actually see the fused layers and all of this let me just stop this. And let me show you that show you right it is going to show you that merging the layers ok.

(Refer Slide Time: 15:40)

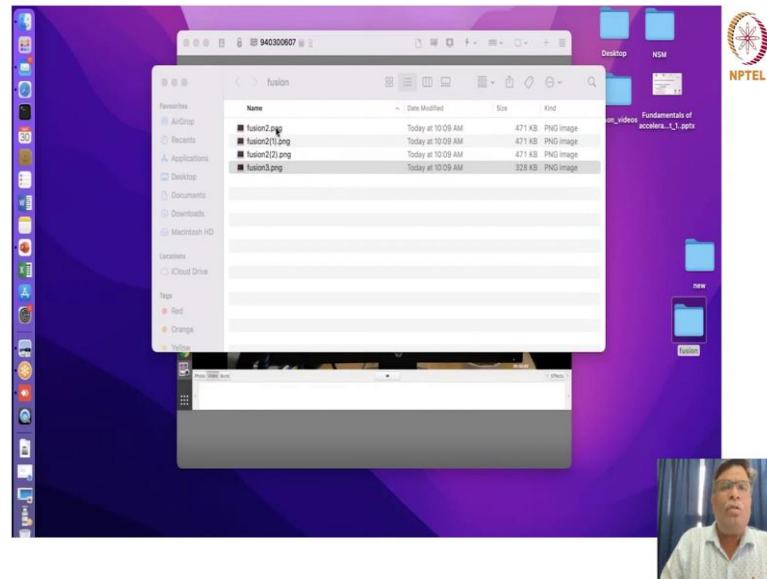


```
root@TK1: /jetson-inference
MobileNetV2/Layer_19_2_Conv2d_5_3x3_s2_128/Relu0
[TRT] After vertical fusions: 94 layers
[TRT] After final dead-layer removal: 94 layers
[TRT] Merging layers: BoxPredictor_0/BoxEncodingPredictor/Conv2D || BoxPredictor_0/ClassPredictor/Conv2D
[TRT] 20
[TRT] Merging layers: BoxPredictor_1/BoxEncodingPredictor/Conv2D || BoxPredictor_1/ClassPredictor/Conv2D
[TRT] 20
[TRT] Merging layers: BoxPredictor_2/BoxEncodingPredictor/Conv2D || BoxPredictor_2/ClassPredictor/Conv2D
[TRT] 20
[TRT] Merging layers: BoxPredictor_3/BoxEncodingPredictor/Conv2D || BoxPredictor_3/ClassPredictor/Conv2D
[TRT] 20
[TRT] Merging layers: BoxPredictor_4/BoxEncodingPredictor/Conv2D || BoxPredictor_4/ClassPredictor/Conv2D
[TRT] 20
[TRT] Merging layers: BoxPredictor_5/BoxEncodingPredictor/Conv2D || BoxPredictor_5/ClassPredictor/Conv2D
[TRT] 20
[TRT] After tensor merging: 88 layers
[TRT] Eliminating concatenation concat_priorbox
[TRT] Generating copy for GridAnchor_1 to concat_priorbox
[TRT] Generating copy for GridAnchor_2 to concat_priorbox
[TRT] Generating copy for GridAnchor_3 to concat_priorbox
[TRT] Generating copy for GridAnchor_4 to concat_priorbox
[TRT] Generating copy for GridAnchor_5 to concat_priorbox
[TRT] After concat removal: 83 layers
[TRT] Graph construction and optimization completed in 0.298781 seconds.
[TRT] Constructing optimization profile number 0 [1/1].
[TRT] ***** Autotuning format combination: -> Float(1,2) *****
[TRT] ***** Autotuning format combination: -> Half(1,2) *****
[TRT] ----- Timing Runner: <reformat> (Reformat)
[TRT] Tactic: 1002 time 0.025469
[TRT] Tactic: 0 time 0.017995
[TRT] Fastest Tactic: 0 time 0.017995
[TRT] ----- Timing Runner: <reformat> (Reformat)
[TRT] Tactic: 1002 time 0.025469
[TRT] Tactic: 0 time 0.018073
[TRT] Fastest Tactic: 0 time 0.018073
[TRT] ***** Autotuning format combination: Float(1,2) -> Float(1,1,4332,8664), Float(1,1,248
0,4800), Float(1,1,600,1200), Float(1,1,226,432), Float(1,1,96,192), Float(1,1,24,48) *****
[TRT] ----- Timing Runner: <reformat> (Reformat)
[TRT] Tactic: 1002 time 0.551432
[TRT] Tactic: 0 time 0.584602
```

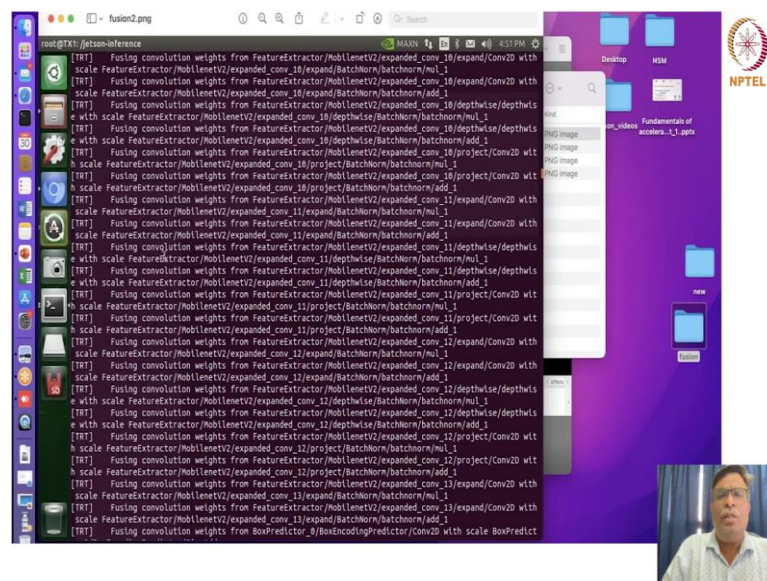
After vertical fusions there are 94 layers after you remove the dead layers right. There 94 layers again ok. And how much auto tuning format combinations right all of this after concatenation removal how many layers are left ok. So, how many layers are left and all of this is going to give you a fair idea of understanding and after you actually merge the tensors ok how many layers are there ok.

So, this basically is trying to give you understanding of how this mobile net program right. Of course, this is a different program from the program which I just now showed right that was for detect net and this I am showing you for mobile net. But that actually also will do the same thing of merging right and fusion and everything of this sort is going to happen.

(Refer Slide Time: 16:48)



(Refer Slide Time: 16:55)



So, in the beginning what is going to happen is you will get something like this right fusing convolution weights from feature extractor for that particular program which particular layer right scale convolution 11 layer. Then fusing convolution layer to then scaling the feature extractor with scaling and everything right.

So, fusing convolution weights and all of this are actually done for lot of things ok. So, this is the next thing. So, it is actually a continuous thing of trying to get all of this right. So, we will try to show you another program which basically is using I do not know it

get stuck again yeah. So, it basically is a detected program ok which you are trying to actually see here the device is a GPU it has loaded things ok and it is a bit slower.

So, it will take time yeah actually the idea is I cannot do SSH onto it because the GUI issue is there. So, I could not have shown you that this basically is a pose detection right which we are trying to do and of course, we can do SSH onto it and I can show you the running command after the GPU thing is over. So, I just wanted to show you this GUI stuff right the window stuff otherwise it would have got stuck.

So, that was the reason I did not use the thing SSH and do it because again it will create a lot of issues on my x server thing. So, that was the reason I had to do something like this. So, I hope it is understandable to some extent in such a short time as to how do you convert ok the program a tensorflow program to a TensorRT program and utilize it on something like a jetson tx 1.