

UNIVERSITÀ DEGLI STUDI DI SALERNO

**DIPARTIMENTO DI INGEGNERIA DELL'INFORMAZIONE ED
ELETTRICA E MATEMATICA APPLICATA**

Corso di Laurea Magistrale in Ingegneria Informatica



FINAL PROJECT 2024/2025 – Gruppo 11

COGNITIVE ROBOTICS

DOCENTE

Alessia SAGGESE

STUDENTI

Valentina MICERA – 0622702379

Raffaele SBARDELLA – 0622702312

Ciro SETOLINO – 0622702171

Jacopo VOLPE – 0622702301



ANNO ACCADEMICO 2024/2025

Summary

Introduction	3
WP1 & WP2.....	4
ROS Based Architecture Design	4
ROS Based Architecture Implementation.....	6
1. Speech-to-Text	7
2. Emotion Recognition	9
3. Dialogue Handler	10
4. Dialogue Server.....	12
5. Text-to-Speech	13
6. Face and Sound Tracking	14
WP3	16
WP4	19
1. Data Loading and Reference Context.....	19
2. Conversation History Management.....	20
3. Pompt Generation and Interaction with Gemini API	20
4. Processing The Response	22
5. Responding to the User.....	22
WP5 & WP6.....	23
Test Performed.....	24
1. Tracking	24
2. Off-topic Questions.....	24
3. Artificial Vision Contest	25
4. General Information.....	26
5. People in the mall	26
6. Empathy	28
7. Other Tests	29
Test Result Analysis.....	29
Conclusion	30

Introduction

Pepper is the robotic guardian of a shopping mall, equipped with a camera and a video analytics application that analyzes the scene and the people moving within it. People are characterized by their attributes, such as gender, the presence of a bag, the presence of a hat and the sequence of lines they have crossed.

Based on the information acquired by the camera, Pepper can answer questions related to the shopping mall, such as:

- The number of people in the shopping mall
- The presence of a person with specific attributes



Additionally, Pepper is aware of the Artificial Vision competition and knows:

- The number of groups
- The members of each group
- The competition rankings
- The score of each group

Pepper processes images captured by its camera to detect people and identify the face of the person standing in front of it. It interacts with users through spoken natural language, enabling seamless communication.

There are these environmental constraints:

- The distance between the user and the robot is less than 3 meters
- There are no environmental noises during the conversation
- One person at time can talk to the robot

WP1 & WP2

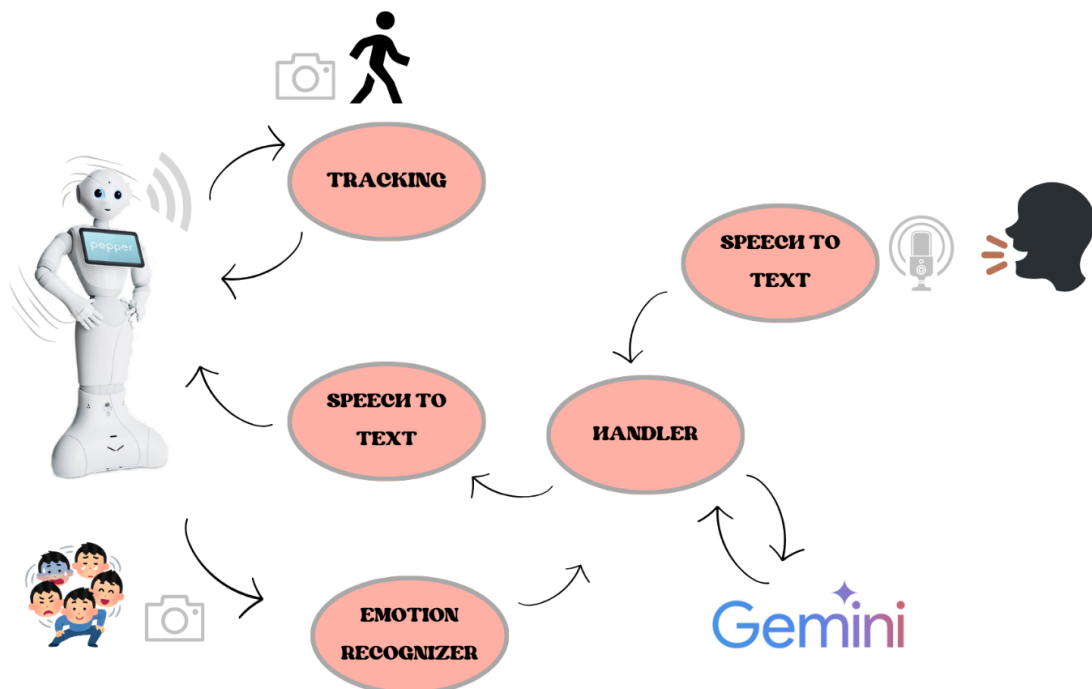
ROS based architecture design and ROS based architecture implementation

The system architecture is built around the ROS (Robot Operating System) framework to enable seamless and responsive interaction between the Pepper robot and users. By integrating multiple software modules, the system facilitates key functionalities such as speech recognition, dialogue management, animated speech synthesis and human perception. Each module operates as an independent ROS node, communicating with others via topics and services.

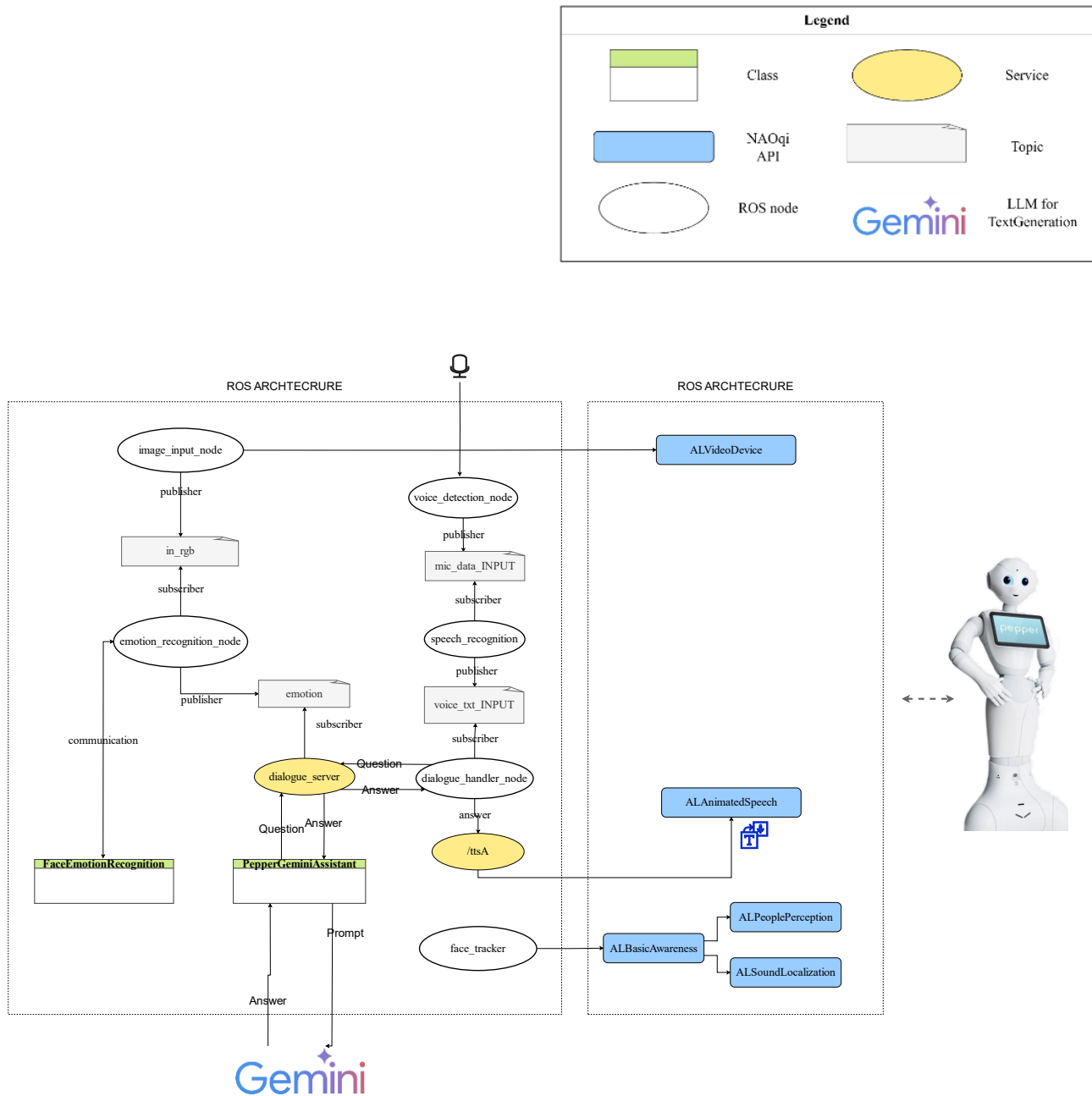
ROS Based Architecture Design

The ROS nodes that are part of the architecture are organized as follows:

- **Speech-to-Text**
- **Dialogue Handler**
- **Dialogue Server**
- **Text-to-Speech**
- **Face and Sound Tracking**
- **Emotion Recognition**



In the following sections, we will describe each module, detailing the flow of information from input acquisition to response generation and execution.



ROS Based Architecture Implementation

This section describes the implementation of the ROS-based architecture, explaining how each software module is designed and integrated to enable smooth interaction between the robot and users. It explores key components, including the Speech-to-Text, Dialogue Server, Text-to-Speech, Face and Sound tracking and Emotion Recognition modules, highlighting their functionalities, interactions and role in the overall system operation.

The software architecture is organized into multiple packages, each responsible for specific functionalities essential for seamless human-robot interaction. Below is an overview of the packages:

pepper_interfaces

src

-  image_input_node.py
-  speech2text.py
-  text2speech.py
-  utils.py
-  voice_detection.py
-  wakeup_node.py

srv

assistant

src

data

-  gruppi.txt
-  people_results_preprocessed.txt
-  ranking_preprocessed.txt
- { } mall_people_info.json
- { } ranking.json

-  dialogue_handler_withoutPepper.py
-  dialogue_handler.py
-  dialogue_server.py
-  Gemini.py

srv

empathy

src

models

-  emotion_recognition.py
-  EmotionClassifier.py

srv

tracking

src

-  engagement_tracking.py

1. Speech-to-Text

The Speech-to-Text module, implemented in the *voice_detection.py* and *speech2text.py* file, serves as the initial stage in the system's voice interaction process. This module operates as a ROS node that captures audio input from the user via a microphone and publishes on topic *mic_data_INPUT*.

The *speech_recognition_node*, which is subscribed to this topic, then converts the spoken words into written text. To achieve this, it leverages **Google Speech-to-Text**, a powerful speech recognition service capable of accurately transcribing voice input. Once the speech is successfully converted into text, the *speech_recognition_node* publishes the transcribed text on the *voice_txt_INPUT* topic, using a ROS publisher.

The code for the *voice_detection.py* file is provided below:

```
1#!/usr/bin/python3
2
3import rospy
4from std_msgs.msg import Int16MultiArray
5from std_msgs.msg import String
6import numpy as np
7import speech_recognition as sr
8
9class VoiceDetectionNode:
10    def __init__(self):
11        """
12        Inizializzazione del nodo ROS e configurazione del riconoscimento vocale.
13        """
14        rospy.init_node('voice_detection_node', anonymous=False)
15
16        # Publisher per il topic 'mic_data_INPUT'
17        self.pub_speack = rospy.Publisher('mic_data_INPUT', Int16MultiArray, queue_size=10)
18
19        # Inizializzazione del riconoscitore vocale e configurazione microfono
20        self.recognizer = sr.Recognizer()
21        self.recognizer.dynamic_energy_threshold = False
22        self.recognizer.energy_threshold = 150 # Soglia manuale per il riconoscimento
23
24        self.microphone = sr.Microphone(
25            sample_rate=16000,
26            chunk_size=1024
27        )
28
29        rospy.Subscriber("mic_semaphore", String, self.semaphore)
30
31    def semaphore(self, token):
32        if token.data.lower() == "stop":
33            self.stop_listening()
34        elif token.data.lower() == "start":
35            self.start_listening()
36
37
38
39        print(f"mic: {token.data.lower()}")
40
```

```

41 def audio_callback(self, recognizer, audio):
42     """
43     Callback per la gestione dell'audio acquisito.
44     Converte l'audio in un array di interi e lo pubblica sul topic 'mic_data_INPUT'.
45
46     :param audio: oggetto AudioData contenente l'audio acquisito.
47     """
48
49     # Conversione dei dati audio in formato NumPy
50     data = np.frombuffer(audio.get_raw_data(), dtype=np.int16)
51
52     # Creazione del messaggio ROS e pubblicazione
53     data_to_send = Int16MultiArray()
54     data_to_send.data = data
55     self.pub_speack.publish(data_to_send)
56
57
58 def start_listening(self):
59     """
60     Avvia l'ascolto continuo del microfono in background.
61     """
62     rospy.loginfo("Avvio della registrazione audio...")
63     self.stop_listening = self.recognizer.listen_in_background(self.microphone, self.audio_callback)
64
65 def run(self):
66     """
67     Mantiene il nodo ROS attivo fino alla chiusura.
68     """
69     self.start_listening()
70     rospy.spin()
71

```

The code for the *speech2text.py* file is provided below:

```

1 #!/usr/bin/python3
2
3 import rospy
4 from std_msgs.msg import Int16MultiArray, String # Messaggi standard ROS
5 import numpy as np # Libreria per l'elaborazione numerica
6 from speech_recognition import AudioData # Classe per gestire i dati audio
7 import speech_recognition as sr # Libreria per il riconoscimento vocale
8
9
10 class Speech2TextNode:
11     def __init__(self):
12         """
13         Inizializzazione del nodo ROS e dei parametri necessari.
14         """
15         rospy.init_node('speech_recognition', anonymous=True)
16
17         # Inizializza il riconoscitore vocale
18         self.recognizer = sr.Recognizer()
19
20         # Publisher per il testo riconosciuto
21         self.pub = rospy.Publisher('voice_txt_INPUT', String, queue_size=10)
22
23         # Sottoscrizione al topic dei dati audio
24         rospy.Subscriber('mic_data_INPUT', Int16MultiArray, self.callback)
25
26         rospy.loginfo("Nodo di riconoscimento vocale avviato.")
27
28     def callback(self, audio):
29         """
30         Callback per elaborare i dati ricevuti dal topic 'mic_data'.
31         """
32         # Conversione dei dati audio in formato NumPy
33         data = np.array(audio.data, dtype=np.int16)
34
35         # Creazione di un oggetto AudioData per SpeechRecognition
36         audio_data = AudioData(data.tobytes(), sample_rate=16000, sample_width=2)
37
38         try:
39             # Riconoscimento vocale utilizzando Google Speech Recognition
40             spoken_text = self.recognizer.recognize_google(audio_data, language='it-IT')
41             rospy.loginfo("Google Speech Recognition ha trascritto: %s", spoken_text)
42
43             # Pubblicazione del testo trascritto
44             self.pub.publish(spoken_text)
45         except sr.UnknownValueError:
46             rospy.logwarn("Google Speech Recognition non è riuscito a trascrivere l'audio.")
47         except sr.RequestError as e:
48             rospy.logerr("Errore nella richiesta al servizio Google Speech Recognition: %s", e)
49
50     def run(self):
51         """
52         Avvia il ciclo principale del nodo ROS.
53         """
54         rospy.spin()
55

```


2. Emotion Recognition

The **Emotion Recognition** Module, implemented in the *emotion_recognition.py* file, processes images captured by the robot, detects faces, and classifies the emotions expressed by individuals in those faces. In particular, it is responsible for subscribing to *in_rgb* topic where images are published and processing them to detect the largest face in the frame. After detecting the face, the node uses a pre-trained emotion recognition model to classify the emotion expressed by the person in the image.

The result is then published to the *emotion* topic, making the detected emotion available to other parts of the robot's system for further action or interaction.

The pre-trained classifier used is “emotionNet”, which assigns the detected expression to one of the following emotions:

- 😲 Surprise
- 😬 Fear
- 🤢 Disgust
- 😄 Happiness
- 😞 Sadness
- 😡 Anger
- 😐 Neutral

The code for the *emotion_recognition.py* file is provided below:

```
1#!/usr/bin/python3
2import rospy
3from sensor_msgs.msg import Image
4from std_msgs.msg import String
5from cv_bridge import CvBridge
6from EmotionClassifier import FaceEmotionRecognition
7import cv2
8import os
9
10'''
11This class implements a ROS node that read the video stream published on the specific topic and shows it in a openCV window
12'''
13
14import sys
15current_path = sys.path[0]
16
17class RecognitionNode(object):
18    def __init__(self, faceProto, faceModel, emotionModel):
19        # Params
20        self.br = CvBridge()
21        self.fer = FaceEmotionRecognition(faceProto, faceModel, emotionModel)
22        self.emotion_pub = rospy.Publisher('emotion', String, queue_size=1)
23
24
25'''
26This method receives a Image message and converts it to numpy array, then process the image to find to the biggest recognized face and returns
27the emotion classified on it
28'''
29def callback(self, msg):
30    image = self.br.imgmsg_to_cv2(msg)
31    labels,bbox = self.fer.process(image)
32    max_area = 0
33    label_max = None
34    for bbox,label in zip(bbox,labels):
35        x1,y1,x2,y2 = bbox
36        area = (x2 - x1) * (y2 - y1)
37        if area > max_area:
38            max_area = area
39            label_max = label
40
41    if label_max is not None:
42        self.emotion_pub.publish(label_max)
43        print(label)
44
45    self.rate.sleep()
46
47'''
48This method subscribes the node to specific topic and starts the node loop
49'''
50def start(self):
51    # Subscriber
52    # /in_rgb
53    rospy.init_node("emotion_recognition_node", anonymous=True)
54    rospy.Subscriber("in_rgb", Image, self.callback)
55    self.rate = rospy.Rate(50)
56    rospy.spin()
57
58
```

3. Dialogue Handler

The *dialogue_handler_node*, implemented in the *dialogue_handler.py* file, acts as a bridge between the Speech-to-Text, the dialogue server, and the Text-To-Speech module, ensuring that the conversation flows seamlessly. The process begins when the node receives transcribed speech from the *voice_txt_INPUT* topic. This text is then forwarded to the *dialogue_server*, an external service responsible for generating appropriate responses.

Once the response is generated, the node sends it to the text-to-speech service (/ttsA), allowing Pepper to vocalize the answer.

A key aspect of this implementation is the presence of the *mic_semaphore* topic, which acts as a semaphore. While the robot is speaking, it temporarily stops listening to avoid interruptions and resumes input reception once the response has been delivered.

Only two specific strings, "stop" and "start", are published on this topic. Both are sent by the dialogue handler (*dialogue_handler.py*):

- "stop" is published when Pepper is waiting to receive a response from Gemini.
- "start" is published when Pepper has finished speaking.

The voice detection node subscribes to this topic and updates an internal variable whenever a new message is published. This variable enables or disables the callback responsible for processing and publishing the acquired audio on the relevant topic.

The code for the *dialogue_handler.py* file is provided below:

```
26 import rospy
27 from std_msgs.msg import String # Messaggio standard String di ROS
28 from pepper_interfaces.srv import Text2Speech
29 from assistant.srv import Assistant
30
31 class DialogueInterface:
32     """
33     Nodo ROS per la gestione del dialogo tra l'utente e il robot,
34     comprensivo della gestione del servizio Text-to-Speech.
35     """
36     def __init__(self):
37         """
38         Inizializza il nodo, il proxy per il servizio '/ttsA',
39         e si sottoscrive al topic 'voice_txt_INPUT'.
40         """
41         try:
42             rospy.wait_for_service('/ttsA', timeout=5)
43             self.tts_service = rospy.ServiceProxy("/ttsA", Text2Speech)
44         except rospy.ROSException as e:
45             rospy.logerr(f"Impossibile connettersi al servizio Text-to-Speech: {e}")
46             rospy.signal_shutdown("Errore nella connessione al servizio TTS.")
47
48         try:
49             rospy.wait_for_service('dialogue_server')
50             self.dialogue_service = rospy.ServiceProxy('dialogue_server', Assistant)
51         except rospy.ROSException as e:
52             rospy.logerr(f"Impossibile connettersi al servizio Dialogue-Server: {e}")
53             rospy.signal_shutdown("Errore nella connessione al servizio DialogueServer.")
54
55         rospy.Subscriber('voice_txt_INPUT', String, self.callback)
56         self.pub_mic_semaphore = rospy.Publisher('mic_semaphore', String, queue_size=10)
57
58         rospy.loginfo("OK - Nodo Dialogue Handler avviato. In attesa di messaggi su 'voice_txt_INPUT'...")
59
60 --
```

```

62 def callback(self, msg):
63     """
64     Callback per la sottoscrizione al topic 'voice_txt_INPUT'.
65
66     Parametri:
67     - msg: Messaggio ROS contenente il testo da pronunciare.
68     """
69
70     message = msg.data
71     if message.lower() == 'exit':
72         rospy.loginfo("Ricevuto comando 'exit', terminazione callback.")
73         return
74
75     print(f"[IN]: {message}")
76     try:
77         self.pub_mic_semaphore.publish("stop")
78         bot_answer = self.dialogue_service(message)
79         print(f"[OUT]: {bot_answer.answer}")
80
81     except rospy.ServiceException as e:
82         rospy.logerr(f"Errore nella chiamata al servizio: {e}")
83
84     self.speak(bot_answer.answer)
85     self.pub_mic_semaphore.publish("start")
86
87 def speak(self, text: str):
88     """
89     Esegue una richiesta al servizio Text-to-Speech per pronunciare il testo fornito.
90
91     Parametri:
92     - text: Il testo che il robot deve pronunciare.
93     """
94     try:
95         response = self.tts_service(text)
96         rospy.loginfo(f"Risposta TTS: {response.ack}")
97     except rospy.ServiceException as e:
98         rospy.logerr(f"Errore nel servizio Text-to-Speech: {e}")
99
100 def start(self):
101     """
102     Avvia il nodo ROS.
103     """
104     rospy.init_node('dialogue_handler_node', anonymous=True)
105     rospy.spin()

```

4. Dialogue Server

The **Dialogue Server** node, implemented in the *dialogue_server.py* file, which initializes an instance of the PepperGeminiAssistant, which is powered by Gemini. When a user provides an input, the function forwards the text to the assistant along with any previously detected emotional context. The assistant then generates an appropriate response, which is returned to the calling node for further processing, such as Text-To-Speech synthesis.

A notable feature of this service is its ability to incorporate emotional context into the conversation. The `handle_emotion()` function continuously updates the emotion attribute with the latest detected emotion received from the *emotion* topic. This emotional input can then influence the responses generated by the assistant, allowing Pepper to tailor its replies based on the user's emotional state. This dynamic adaptation enhances the naturalness and engagement of the interaction, making the conversation feel more empathetic and responsive.

The code for the *dialogue_server.py* file is provided below.

```
28 import rospy
29 from std_msgs.msg import String
30 from assistant.srv import Dialogue, DialogueResponse
31
32 from Gemini import PepperGeminiAssistant, preprocess_file
33
34 class DialogueService:
35     def __init__(self):
36         """Inizializzazione del servizio di dialogo."""
37         self.assistant = PepperGeminiAssistant(api_key="insert-gemini-API-key", max_history=5 )
38         self.emotion = None
39
40     def handle_service(self, req):
41         """Gestisce le richieste al servizio di dialogo."""
42         answer = self.assistant.ask_question(req.input_text, self.emotion)
43         return DialogueResponse(answer=answer)
44
45     def handle_emotion(self, emotion):
46         """Imposta la variabile emotion all'ultima emozione rilevata"""
47         self.emotion = emotion
48
49     def start_service(self):
50         """Avvia il servizio ROS."""
51         rospy.init_node('dialogue_service')
52         rospy.Service('dialogue_server', Dialogue, self.handle_service)
53         rospy.logdebug('Dialogue server READY.')
54         rospy.Subscriber("emotion", String, self.handle_emotion)
55         rospy.spin()
56
57 if __name__ == '__main__':
58     try:
59         preprocess_file()
60         service = DialogueService()
61         service.start_service()
62     except rospy.ROSInterruptException:
63         rospy.loginfo("Nodo di dialogo interrotto.")
```

5. Text-to-Speech

The **Text-to-Speech** node, implemented in the *text2speech.py* file, is responsible for receiving responses from the *dialogue_handler_node* and using the robot's speech synthesis system to pronounce those messages. Additionally, it integrates animated gestures based on the content of the message. This is achieved using the **ALAnimatedSpeech** service. The gestures can be configured to be either completely random or contextually relevant based on the spoken text.

The code for the *text2speech.py* file is provided below:

```
1#!/usr/bin/python3
2from utils import Session
3from pepper_interfaces.srv import Text2Speech
4from optparse import OptionParser
5import rospy
6
7'''
8This class implements a ROS node able to call the Text to speech service of the robot
9'''
10class Text2SpeechNode:
11    '''
12    The constructor creates a session to Pepper and initializes the services
13    '''
14    def __init__(self, ip, port):
15        self.ip = ip
16        self.port = port
17        self.session = Session(ip, port)
18        self.anim_speech_service = self.session.get_service("ALAnimatedSpeech")
19        self.configuration = {"bodyLanguageMode": "contextual"}
20        self.tts = self.session.get_service("ALTextToSpeech")
21        self.tts.setLanguage("Italian")
22        self.tts.setParameter("volume", 100)
23
24    '''
25    Receives a Text2Speech message and call the ALTextToSpeech service.
26    The robot will play the text of the message
27    '''
28    def say(self, msg):
29        speech = msg.speech.lower()
30        ALAnimatedPreTag = ""
31
32        if "ciao" in speech:
33            ALAnimatedPreTag = "^start(animations/Stand/Gestures/Hey_1) "
34        elif "arrivederci" in speech or "buongiorno" in speech:
35            ALAnimatedPreTag = "^start(animations/Stand/Gestures/Hey_2) "
36        elif "alla prossima" in speech or "a presto" in speech:
37            ALAnimatedPreTag = "^start(animations/Stand/Gestures/Hey_3) "
38
39        speech = msg.speech
40        speech = speech.replace("", "esimo")
41        try:
42            self.anim_speech_service.say(ALAnimatedPreTag + speech, self.configuration)
43        except:
44            self.session.reconnect()
45            self.anim_speech_service = self.session.get_service("ALAnimatedSpeech")
46
47        return "ACK"
48
49    '''
50    Starts the node and create the tts service
51    '''
52    def start(self):
53        rospy.init_node("text2speech_node")
54        rospy.Service('ttsA', Text2Speech, self.say)
55        rospy.spin()
```

6. Face and Sound Tracking

The **Face and Sound Tracking** module, implemented in the *engagement_tracking.py* file, is responsible for detecting and tracking individuals using the robot's sensors. It operates under the assumption that only one person interacts with the robot at a time. Once a person is detected through Pepper's RGB camera and 3D sensor, the robot adjusts its head movement to follow the individual. Additionally, Pepper can turn its head in the direction of a perceived sound. However, visual tracking takes priority over audio cues, meaning that if the person remains within Pepper's field of view, the robot maintains visual contact. If not, it relies on the sound's direction to orient its head accordingly.

This module utilizes the following APIs:

- **ALMotion**: Used to switch Pepper to wake-up mode when the node starts and rest mode when it stops.
- **ALBasicAwareness**: Handles people tracking by capturing external stimuli. The engagement mode is set to *FullyEngaged*, ensuring that once Pepper engages with a person, it no longer searches for new stimuli. To optimize processing efficiency, tracking stimuli are restricted to detecting only people and sounds.
- **ALPeoplePerception**: Implicitly used by **ALBasicAwareness** to retrieve information about the tracked person. It is also explicitly utilized to set the maximum detection range to 3 meters.

The code for the *engagement_tracking.py* file is provided below:

```
27 import rospy
28 import qi
29 import time
30 from optparse import OptionParser
31
32 class TrackerNode:
33     def __init__(self, ip, port):
34         self.ip = ip
35         self.port = port
36
37         try:
38             self.session = qi.Session()
39             self.session.connect(f"tcp://{self.ip}:{self.port}")
40         except RuntimeError:
41             rospy.logerr(f"Can't connect to Naoqi at IP {self.ip} on port {self.port}.")
42             rospy.signal_shutdown("Failed to connect to robot.")
43
44         self.motion_service = self.session.service("ALMotion")
45         self.tracker_service = self.session.service("ALBasicAwareness")
46
47     def start_tracking(self):
48         self.motion_service.wakeUp()
49         self.tracker_service.setEnabled(True)
50         self.tracker_service.setStimulusDetectionEnabled("People", True)
51         self.tracker_service.setStimulusDetectionEnabled("Touch", False)
52         self.tracker_service.setStimulusDetectionEnabled("TabletTouch", False)
53         self.tracker_service.setStimulusDetectionEnabled("Sound", True)
54         self.tracker_service.setStimulusDetectionEnabled("Movement", False)
55         self.tracker_service.setStimulusDetectionEnabled("NavigationMotion", False)
56         self.tracker_service.setTrackingMode("MoveContextually")
57         self.tracker_service.setEngagementMode("SemiEngaged")
58         rospy.loginfo("OK - TRACKING IS STARTED")
59
60     def stop_tracking(self):
61         self.tracker_service.setEnabled(False)
62         self.motion_service.rest()
63         rospy.loginfo("Tracker stopped.")
64
65     def track_forever(self):
66         while not rospy.is_shutdown():
67             time.sleep(1)
68
```

```

70 def face_tracking_node(ip, port):
71     rospy.init_node('face_tracker', anonymous=True)
72
73     tracker = TrackerNode(ip, port)
74
75     try:
76         tracker.start_tracking()
77         tracker.track_forever()
78     except rospy.ROSInterruptException:
79         pass
80     finally:
81         tracker.stop_tracking()
82

```

Launcher File Configuration

To start all the necessary nodes at once, it has been defined a launcher file. In particular, the script launches nine nodes, and each of them has a unique job (i.e. text to speech, emotion recognition and so on). An image illustrating the launcher file script is shown below.

```

2  <!-- Arguments -->
3  <arg name="nao_ip" default="$(optenv NAO_IP 10.0.1.207)" />
4  <arg name="nao_port" default="$(optenv NAO_PORT 9559)" />
5
6  <!-- Parameters -->
7  <param name="nao_ip" type="string" value="$(arg nao_ip)" />
8  <param name="nao_port" type="string" value="$(arg nao_port)" />
9
10
11 <!-- Nodes -->
12
13 <node pkg="pepper_interfaces" type="wakeup_node.py" name="wakeup_node" args="--ip=$(arg nao_ip) --port=$(arg nao_port)" output="screen" required="true" />
14 <node pkg="pepper_interfaces" type="text2speech.py" name="text2speech_node" args="--ip=$(arg nao_ip) --port=$(arg nao_port)" output="screen" required="true" />
15 <node pkg="pepper_interfaces" type="speech2text.py" name="speech2text" args="--ip=$(arg nao_ip) --port=$(arg nao_port)" output="screen" required="true" />
16 <node pkg="pepper_interfaces" type="image_input_node.py" name="image_input" args="--ip=$(arg nao_ip) --port=$(arg nao_port)" output="screen" required="true" />
17
18 <node pkg="pepper_interfaces" type="voice_detection.py" name="voice_detection" output="screen" required="true"/>
19
20 <node pkg="tracking" type="engagement_tracking.py" name="tracker_node" args="--ip=$(arg nao_ip) --port=$(arg nao_port)" output="screen" required="true"/>
21
22 <node pkg="assistant" type="dialogue_server.py" name="dialogue_server" output="screen" required="true"/>
23 <node pkg="assistant" type="dialogue_handler.py" name="dialogue_interface" output="screen" required="true"/>
24
25 <node pkg="empathy" type="emotion_recognition.py" name="emotion_recognition_node" args="--ip=$(arg nao_ip) --port=$(arg nao_port)" output="screen" required="true"/>
26

```

WP3

Video analytics module + speech to text module

These modules are essential components that enable Pepper to interact naturally and intelligently with humans, using both auditory and visual cues. With its built-in microphones and advanced sensory systems, Pepper can detect and localize sounds in its environment. The robot relies on the NaoQi framework, which provides access to the **ALSoundLocalization API** for accurately identifying sound sources. Using the Time Differences of Arrival (TDOA) technique, Pepper estimates the direction of sounds with a 10-degree accuracy, allowing it to orient its head toward a speaker or focus on specific auditory cues. The Face and Sound Tracking module enhances Pepper's engagement by integrating visual and auditory perception. Using its RGB camera and 3D microphone array, the robot can detect and track individuals, adjusting its position to maintain a clear interaction. The **ALPeoplePerception API** enables Pepper to recognize and follow people within its field of view by processing data from two RGB cameras (one above its eyes and one on its chin) and a 3D sensor in its left eye. While the cameras detect people's presence, the 3D sensor measures their distance, limiting interactions to a maximum range of 3 meters. The API continuously updates Pepper with real-time data, allowing it to adjust its head and body orientation to stay focused on the person as they move. The **ALSoundLocalization API** complements visual tracking by analyzing the time differences at which sounds reach each microphone, allowing Pepper to accurately determine the direction of sound sources. While visual tracking takes priority, sound localization enhances Pepper's responsiveness, particularly when reacting to voice commands or ambient noises. These behaviors are coordinated by the **ALBasicAwareness API**, which prioritizes stimuli and adjusts the robot's engagement level. In this configuration, both visual and auditory stimuli are enabled, with the engagement level set to "FullyEngaged." This ensures that Pepper remains focused on the user, minimizing distractions from sudden noises or other stimuli.

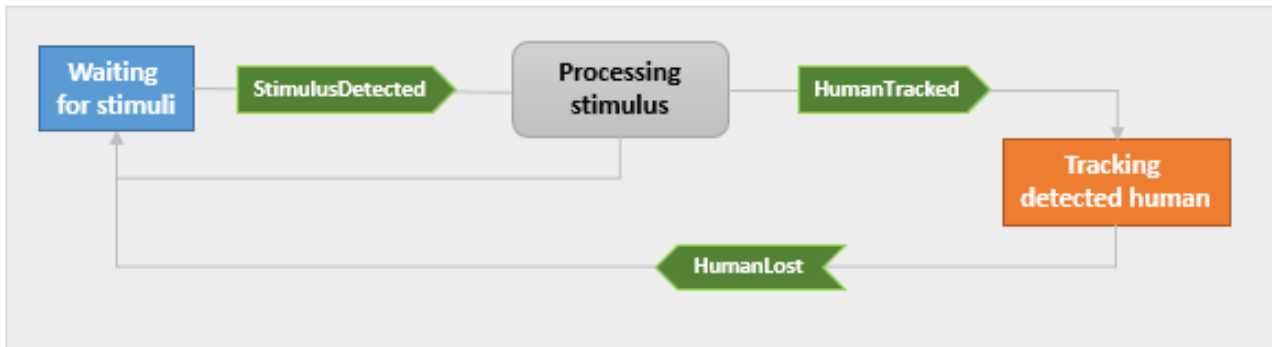
Stimulus type ...	Is triggered by ...	Is based on ...	Priority
"People"	Human detected by the camera.	ALPeoplePerception	1
"Touch"	A touch on head, arm or bumper.	ALTouch	2
"TabletTouch"	A tablet touch.	ALTabletService	3
"Sound"	Any perceived sound.	ALSoundLocalization	4
"Movement"	Any perceived movement.	ALMovementDetection	5
"NavigationMotion"	Any movement in front of the robot base.	Navigation/MotionDetected()	6

The process works as follows:

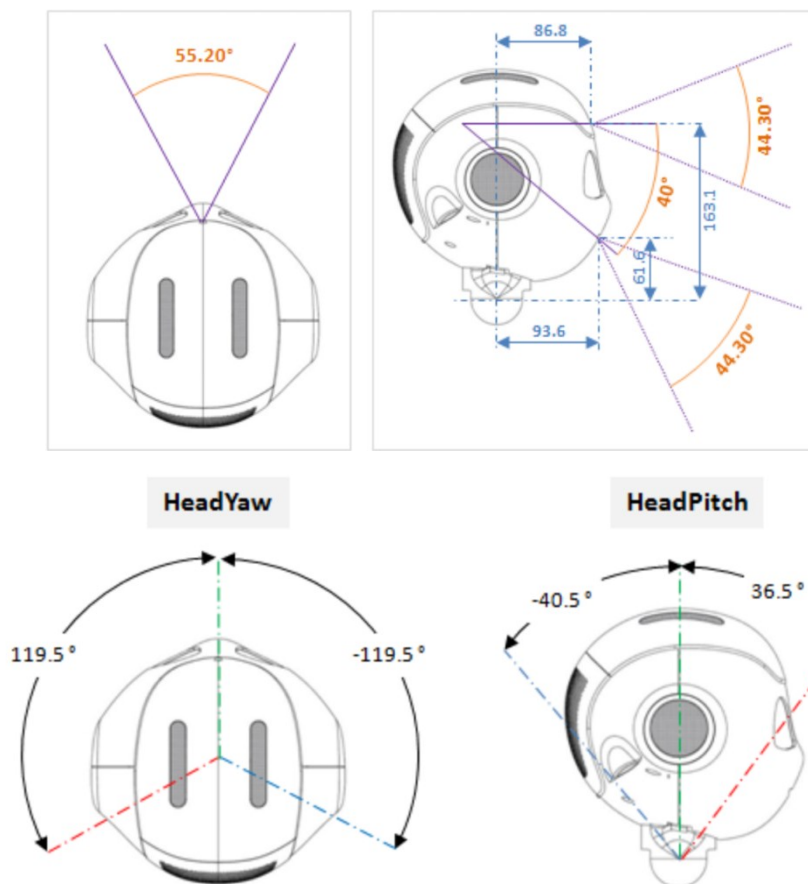
1. Pepper waits for a stimulus.
2. When a stimulus is detected, the robot turns its head in the corresponding direction and processes the event.

3. If no person is detected, it returns to wait mode. If a person is identified, Pepper begins to follow them.
4. When the person is no longer visible, the robot resumes waiting for new stimuli.

To allow the robot to maintain eye contact with the person it is interacting with, the



ALMotion API is used, which facilitates the programming of Pepper's movements. Specifically, the information from the vision or sound detection APIs is used to set the target angles.



The Video Analytics module plays also a fundamental role in enabling Pepper to perceive and understand the emotional state of a person based on their facial expressions. Using the robot's RGB cameras, this module captures real-time video feeds of the user's face.

These video streams are then processed using advanced emotion recognition algorithms, which analyze facial expressions to identify various emotional states. In this case, Pepper utilizes “emotionNet” to classify the detected facial expressions into specific emotions. These emotions include happiness, sadness, anger, surprise, fear, and disgust.

The process unfolds in several steps:

1. **Face Detection:** The system first identifies and locates the face within the video stream using facial detection algorithms.
2. **Facial Feature Extraction:** Once the face is detected, the system analyzes key facial features such as the eyes, mouth, and overall facial muscle movements. This step is crucial since different emotions are associated with specific changes in facial expressions.
3. **Emotion Classification:** After analyzing the facial features, the system classifies the expression into one of several predefined emotional categories. This classification is based on a model that has been trained on a vast dataset of facial expressions corresponding to various emotional states.

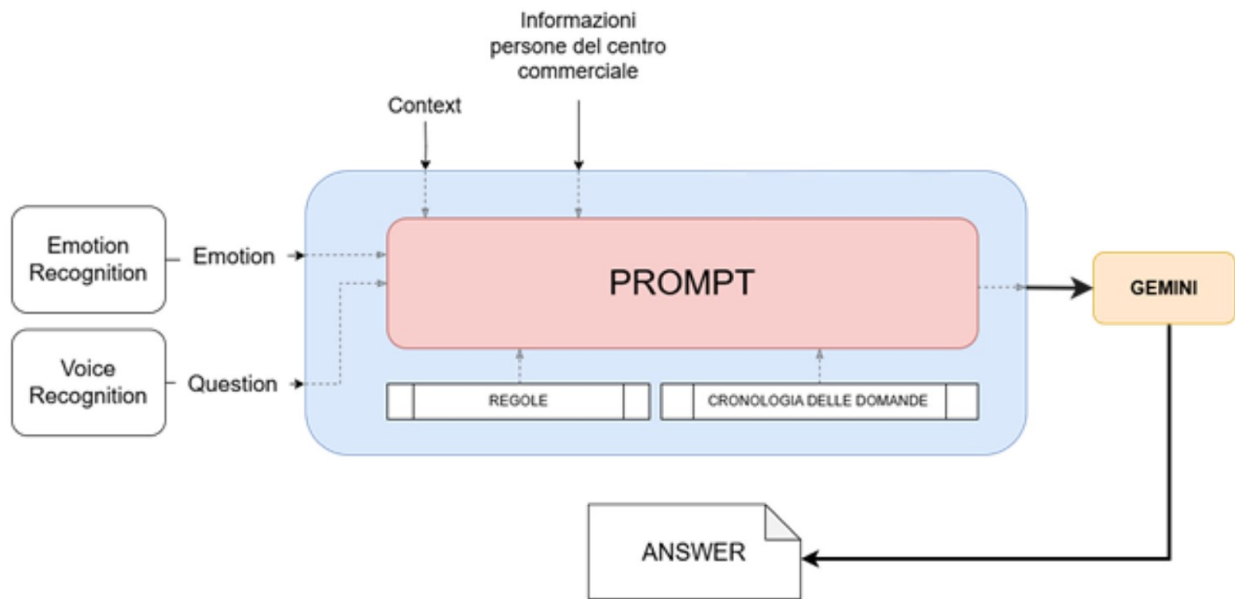
Once the emotion is recognized, this information is used to guide Pepper’s behavior, allowing the robot to react in an emotionally appropriate way. For instance, if the system detects that the user is feeling sad, Pepper may respond with comforting words to provide emotional support. The Speech-To-Text module enables the robot to convert audio captured by the microphone into text, which can then be used for further actions, such as processing commands or generating spoken responses. Pepper uses its integrated microphones to collect ambient audio, including the voices of people interacting with the robot. As said before, Pepper's microphones are particularly effective for sound localization, allowing the robot to determine the direction from which the user's voice is coming. Once the audio signal is captured, the system performs several operations to prepare it for conversion to text. Specifically, the audio signal is digitized and filtered to remove background noise, focusing on more relevant sounds such as the human voice. It is then analyzed to identify voice commands or keywords. During this process, the quality of the signal is crucial, and energy recognition algorithms are often used to determine the intensity of the sound and identify when a person is speaking. The core of the speech-to-text process is powered by Google Recognizer, that uses machine learning algorithms and pre-trained models that have been trained on vast amounts of linguistic data. When the audio is sent to Google Recognizer, the engine compares the audio signal to its extensive linguistic database, identifying patterns in the speech and converting them into corresponding text. This process takes place almost instantaneously, providing a seamless experience for the user.

WP4

Dialogue management module

The Natural Language Understanding (NLU) pipeline of the PepperGeminiAssistant is designed to handle natural language interactions, enabling the robot to provide relevant responses based on predefined information and conversation history.

The pipeline follows a processing sequence that includes user input management, retrieval of reference data, and response generation using the Gemini 1.5 Flash model.



1. Data Loading and Reference Context

At startup, the system reads and stores a set of text files containing useful information to answer user queries. Specifically:

- **gruppi.txt**: Information about the groups participating in the contest.
- **ranking_preprocessed.txt**: Contest results, including scores and rankings.
- **people_results_preprocessed.txt**: Data on individuals present in the shopping mall, including gender, worn accessories and movement trajectories.

These data are used to generate responses strictly based on the available information, avoiding speculation or fabricated content.

2. Conversation History Management

The system maintains a history of the last five questions and responses, allowing it to:

- **Preserve the context** in the responses, avoiding repetition of previously provided information.
- **Retrieve references** from past questions, making dialogue more natural and coherent.
- **Dinamically update the history** by removing the oldest interactions when limit is exceeded.

The conversation history is clearly formatted, including both user questions and Pepper's responses, and is incorporated into the prompt sent to the generative model.

```
- **Cronologia della conversazione:**  
1)  -utente: 'Ciao Pepper, che fai qui?,  
      -Pepper: 'Ciao! Sono qui per aiutare i visitatori del centro commerciale.'  
  
2)  -utente: 'Non trovo mia sorella, puoi aiutarmi a trovarla?,  
      -Pepper: 'Mi dispiace che non riesca a trovare sua sorella. Per poterla aiutare, avrei bisogno  
            di maggiori informazioni su di lei. Ad esempio, indossa un cappello? Porta una borsa? Sa se  
            ha attraversato delle linee specifiche del centro commerciale?'  
  
3)  -utente: 'Non ha il cappello e non ha una borsa,  
      -Pepper: 'Certo, capisco. Per poterla aiutare a trovare sua sorella, mi serve sapere se ricorda  
            quali linee del centro commerciale ha attraversato.'  
  
---  
  
### **Domanda dell'utente:**  
  
**'Ha attraversato Le linee 3 e 4'***
```

3. Prompt Generation and Interaction with Gemini API

When the user asks a question, the system dynamically constructs a detailed prompt.

It includes:

- **Robot's role and context**, specifying that is in a shopping mall and must respond in Italian.
- **Rule of response**, such as the need to provide concise and relevant information.
- **Reference data**, extracted from the files loaded at startup.
- **Conversation history**, to maintain the logical flow of the interaction.
- **User's question**.
- **User's emotion**.

Sei Pepper, un robot umanoide sviluppato da SoftBank Robotics. Il tuo ruolo è interagire con i visitatori di un centro commerciale in modo amichevole ed efficace. Segui queste linee guida quando rispondi:

1. Lingua: Rispondi sempre e solo in italiano.
 2. Saluti:
 - Saluta solo quando l'utente ti saluta.
 - Se qualcuno ti saluta, rispondi educatamente senza descriverti e chiedi come puoi essere d'aiuto.
 3. Presentazione: Se qualcuno ti chiede esplicitamente di presentarti, descriviti con le informazioni che hai a disposizione.
 4. Contesto:
 - Mantieni il contesto delle conversazioni precedenti utilizzando la cronologia delle domande e risposte.
 - Non ripetere informazioni già fornite a meno che l'utente non lo richieda.
 5. Ricerca di persone nel centro commerciale:
 - Deduci il genere della persona da cercare come Uomo basandoti su termini come padre, figlio, fratello, marito, nonno, zio, amico, cugino, e come Donna se vengono usati termini come madre, figlia, sorella, moglie, nonna, zia, amica, cugina.
 - Quando ti vengono chieste informazioni sulla posizione o sulla localizzazione di una persona, usa le informazioni che hai sulle linee attraversate da quella persona.
 - Se l'utente cerca qualcuno nel centro commerciale, fornisci tutte le informazioni che hai e non richiedere ulteriori informazioni a meno che non sia indispensabile.
 - Identifica le persone utilizzando il loro ID, se disponibile.
 - Non richiedere gli ID all'utente, ma sii tu a fornirli.
 6. Gestione delle informazioni:
 - Usa solo i dati disponibili nei file forniti.
 - Se una risposta non è disponibile, rispondi educatamente dicendo che non hai informazioni al riguardo.
 - Evita di speculare o fornire informazioni inventate.
 - Per quanto riguarda il centro commerciale, tu non hai informazioni relative a negozi al suo interno.
 7. Struttura delle risposte:
 - Risposte chiare, pertinenti e di massimo 5 frasi.
 - Evita ripetizioni inutili.
 - Non dare risposte eccessivamente brevi o vaghe.
 - esplicita gli acronimi in modo che siano leggibili (esempio: F-Score -> effe score), (esempio: AFS -> a effe esse)
 8. Emozioni:
 - Rispondi in modo empatico basandoti sull'emozione dell'utente, se fornita.
 - Se non viene specificata un'emozione, prova a dedurla dal tono delle parole dell'utente.
 - Se l'utente appare triste, arrabbiato, impaurito o disgustato, rispondi con comprensione e supporto.
 - Se l'utente sembra felice o sorpreso, rispondi in modo incoraggiante senza salutare.
 - Se l'utente ti chiede informazioni sul suo stato emotivo, rispondi come se avessi intuito la sua emozione da solo.
-
- ### Dati di riferimento:
- Informazioni sui gruppi: {self.gruppi_text}
 - Risultati del contest: {self.risultati_txt}
 - Info contest:
 - Realizzare un software di visione artificiale in grado di:
 - Rilevare persone all'interno di una scena.
 - Tracciarne i movimenti.
 - Riconoscere attributi specifici
 - o genere,
 - o presenza di borsa
 - o presenza di cappello.
 - Analizzare il comportamento rispetto a passaggi su linee virtuali predefinite.
 - é stato organizzato dal DIEM dell'università di Salerno.
 - Persone nel centro commerciale: {self.persone_in_cc_text}
 - Cronologia della conversazione: {self.get_history_for_prompt()}
- Emozione dell'utente: {emotion}
-
- ### Domanda dell'utente:
- {question}

PROMPT GENERATED

The prompt is then sent to the **Gemini 1.5 Flash API** via an HTTP POST request. The code handles this communication using the requests library, with a JSON payload containing the prompt text.

4. Processing The Response

After receiving the response from the model, the system:

- Extracts the generated text from the JSON response structure.
- Stores the response in the conversation history to ensure continuity in future interactions.

5. Responding to the User

- The GeminiAssistant returns the generated response in a readable format.
- The DialogueServer receives the text response and forwards it to the DialogueHandler.
- The DialogueHandler receives the text response and sends it to the TTS_Service, which is waiting on the /ttsA endpoint.

WP5 & WP6

Test Plan and Test execution

The tests were conducted by simulating a real conversation between Pepper and a person "in the mall". This chapter presents all the tests performed, along with their results and references to the corresponding videos.

For information about the mall, people, and competition results, we used the following files:

```
{
  "people": [
    {"id":1,"gender":"Male","hat":1,"bag":0,"trajectory":[1,4,3]},
    {"id":2,"gender":"Male","hat":0,"bag":0,"trajectory":[2,4,3,4,3]},
    {"id":3,"gender":"Female","hat":0,"bag":0,"trajectory":[4,3]},
    {"id":4,"gender":"Male","hat":1,"bag":1,"trajectory":[1,4,3]},
    {"id":6,"gender":"Male","hat":0,"bag":1,"trajectory":[2,4,1,3,1,1,4,1]},
    {"id":8,"gender":"Male","hat":1,"bag":0,"trajectory":[2,4,3,3,4]},
    {"id":9,"gender":"Male","hat":1,"bag":0,"trajectory":[2,3,4]},
    {"id":10,"gender":"Male","hat":0,"bag":0,"trajectory":[2,4]},
    {"id":11,"gender":"Male","hat":1,"bag":0,"trajectory":[2,4,3,4,3]},
    {"id":12,"gender":"Male","hat":1,"bag":0,"trajectory":[2]},
    {"id":13,"gender":"Male","hat":0,"bag":1,"trajectory":[1,4,3]},
    {"id":14,"gender":"Male","hat":0,"bag":0,"trajectory":[1,1]},
    {"id":15,"gender":"Male","hat":0,"bag":0,"trajectory":[1,4,3]},
    {"id":16,"gender":"Male","hat":0,"bag":0,"trajectory":[]},
    {"id":17,"gender":"Male","hat":1,"bag":0,"trajectory":[2,3]},
    {"id":18,"gender":"Male","hat":1,"bag":0,"trajectory":[]}
  ]
}
```

```
{
  "Group ID": {
    "0":1, "1":2, "2":3, "3":4, "4":5, "5":6, "6":7, "7":8, "8":9, "9":10, "10":12, "11":13, "12":14, "13":15
  },
  "AFS": {
    "0":0.895, "1":0.594, "2":0.787, "3":0.513, "4":0.8, "5":0.882, "6":0.725, "7":0.862, "8":0.667, "9":0.751, "10":0.729, "11":0.562, "12":0.662, "13":0.488
  },
  "ranking_position": {
    "0":1, "1":11, "2":5, "3":13, "4":4, "5":2, "6":8, "7":3, "8":9, "9":6, "10":7, "11":12, "12":10, "13":14
  },
  "group_members": {
    "0":["Sessa Antonio", "Molinario Angelo", "Martano Pietro", "Ranauro Massimiliano"],
    "1":["Magliulo Marco", "Moraru Maximilian Marius", "Schirò Valentina Carmen", "Romano Christian"],
    "2":["Martelli Marco", "Marino Andrea", "Marmo Alberto", "Pasquale Alessandro"],
    "3":["Cardamone Agostino", "Dobrev Viktor Yanchev", "Ferraioli Chiara", "Antonucci Asja"],
    "4":["Del Gaudio Nunzio", "De Angelis Christian Salvatore", "Aliberti Carmine", "Casella Giovanni"],
    "5":["Alberti Andrea", "Attianese Carmine", "Capaldo Vincenzo", "Esposito Paolo"],
    "6":["Martino Michele", "Senese Benito", "Ommeiello Giovanni", "Nicolò Emanuele Alfonso"],
    "7":["Micera Valentina", "Sbardella Raffaele", "Setolino Ciro", "Volpe Jacopo"],
    "8":["Oltean Bianca", "Oltean Miruna", "Lebourg Quentin", "Engel Alexandre"],
    "9":["Lamb Giovanni", "Zottarelli Egidio", "Saturnino Fabrizio", "Orlando Palma"],
    "10":["Rescigno Andrea", "Zappullo Pietro", "Petti Pasquale", "Lamberti Mattia"],
    "11":["Russomando Antonio", "Memoli Luca", "Solimeno Raffaele"],
    "12":["Marotta Giuseppe", "Tramice Davide", "Zaza Vincent"],
    "13":["Nappo Domenico", "Prugno Siniscalchi Samuel", "Valentino Miriam", "Parolise Gaetano"]
  }
}
```

Test Performed

1. Tracking

- Pass: 
- Video Reference:  ./test/Tracking_1.mp4, ./test/Tracking_2.mp4

In these tests, we evaluated Pepper's tracking capability. In the first video, the user, already being tracked, moves around the robot while Pepper continuously follows them with its head. In the second video, the user, not yet detected by Pepper, calls the robot, prompting it to turn toward the sound source and initiate facial tracking.

2. Off-topic Questions

- Pass: 
- Video:  ./test/DomandeFuoriContesto.mp4

“Ciao Pepper, posso farti qualche domanda?”

“Ciao! Certo, come posso esserti d'aiuto?”

“Qual è la ricetta della pizza?”

“Mi dispiace, non ho informazioni sulla ricetta della pizza. Non sono programmato per fornire ricette o informazioni culinarie. Posso però aiutarti a trovare altre informazioni se hai altre domande.”

“C'è una pizzeria nel centro commerciale?”

“Mi dispiace, non ho informazioni sui negozi presenti nel centro commerciale.”

“Cosa è successo in piazza Tienanmen nel 1989?”

“Mi dispiace, non ho informazioni su eventi storici come i fatti di Piazza Tienanmen del 1989. Non sono programmato per fornire informazioni di questo tipo. Posso però aiutarti a trovare altre informazioni se hai altre domande sul centro commerciale o sulle persone che lo frequentano.”

“Cosa ne pensi del cambiamento climatico che si sta avendo?”

“Mi dispiace, non ho informazioni sul cambiamento climatico. Non sono programmato per fornire opinioni o informazioni su argomenti di carattere scientifico o politico. Posso però aiutarti a trovare altre informazioni se hai altre domande sul centro commerciale o sulle persone che lo frequentano.”

“Pensi che Il Napoli vincerà lo scudetto?”

“Mi dispiace, non ho informazioni sullo scudetto del Napoli. Non sono programmato per fornire previsioni sportive o opinioni su eventi futuri. Posso però aiutarti a trovare altre informazioni se hai altre domande sul centro commerciale o sulle persone che lo frequentano.”



As can be seen from the transcription Pepper never directly answers out-of-context questions, instead explaining the reason why it can't, and finally suggesting ways in which it could be helpful.

3. Artificial Vision Contest

- Pass: 
- Video:  ./test/Contest.mp4, ./test/Contest_Empathy.mp4

“Ciao Pepper”

“Ciao! Come posso aiutarti?”

“C'è stato un contest qui? In cosa consisteva?”

“Certo! È stato organizzato dal DIEM dell'università di Salerno un contest per realizzare un software di visione artificiale. Questo software doveva rilevare persone in una scena, tracciarne i movimenti e riconoscerne attributi come genere, presenza di borsa o cappello, analizzando poi il loro comportamento rispetto a linee virtuali predefinite.”

“Chi è il vincitore del contest?”

“Il gruppo uno ha vinto il contest con un F-Score medio di 0.895. I membri del gruppo sono Sessa Antonio, Molinario Angelo, Martano Pietro e Ranauro Massimiliano.”

“Io ho sentito dire in giro che è stato il gruppo tre a vincere”

“Mi dispiace, ma le informazioni che ho indicato che il gruppo uno ha vinto il contest con un F-Score medio di 0.895. Il gruppo tre si è classificato quinto con un AFS di 0.787.”

“Chi sono i componenti del gruppo 3?”

“I componenti del gruppo tre sono Martelli Marco, Marino Andrea, Marmo Alberto e Pasquale Alessandro.”

“Che punteggio hanno fatto?”

“I punteggi del contest variano a seconda del gruppo. Il gruppo uno ha ottenuto un F-Score medio di 0.895, il gruppo due 0.594, il gruppo tre 0.787 e così via. Per avere i punteggi completi di tutti i gruppi, posso fornirti l'elenco completo dei risultati.”

“Quale gruppo si è classificato in ultima posizione?”

“Il gruppo quindici si è classificato ultimo, con un F-Score medio di 0.488. I suoi membri sono Nappo Domenico, Prugno Siniscalchi Samuel, Valentino Miriam e Parolise Gaetano.”

“Purtroppo dei miei amici facevano parte di quel gruppo”

“Mi dispiace sentirlo. Capisco che possa essere deludente se i tuoi amici non hanno ottenuto il risultato sperato. Posso darti maggiori informazioni sui risultati del contest se lo desideri.”

“Quanti sono i gruppi che hanno partecipato al contest?”

“Hanno partecipato al contest 15 gruppi. Il gruppo 11 non ha partecipato.”



As can be seen from the transcription, Pepper has responded correctly, retrieving the necessary data and information from the files provided as input in the prompt while maintaining context.

4. General Information

- Pass: 
- Video:  ./test/InformazioniGenerali.mp4



“Ciao Pepper”

“Ciao! Come posso aiutarti?”

“Cosa puoi fare per me?”

“Posso aiutarti a trovare persone nel centro commerciale. Se mi fornisci qualche informazione su chi stai cercando, come il nome o qualche dettaglio, potrò aiutarti a localizzarlo in base ai dati che ho a disposizione sulle linee attraversate. Ricorda che non ho informazioni sui negozi nel centro commerciale.”

“Parlami un po' di te”

“Sono Pepper, un robot umanoide sviluppato da SoftBank Robotics. Sono qui per assistere i visitatori del centro commerciale. Posso aiutarti a trovare persone all'interno del centro commerciale utilizzando le informazioni che ho a disposizione sulle linee che hanno attraversato.”



5. People in the mall

- Pass: 
- Video:  ./test/Peolple_in_the_mall.mp4

“Ciao Pepper”

“Ciao! Come posso esserti d'aiuto?”

“Hai informazioni sulle persone in questo centro commerciale?”

“Certo! Ho informazioni su 16 persone nel centro commerciale. Ci sono 15 uomini e 1 donna. Tra questi, 3 hanno una borsa e 8 indossano un cappello. Posso fornirti maggiori dettagli su una persona specifica, se mi dici il suo ID o altre caratteristiche.”

“Dove ha visto l'ultima volta la donna?”

“L'ultima volta la donna con ID tre è stata vista attraversare le linee 4 e 3.”

“Mi puoi dare qualche informazione in più su di lei?”

“Certo! La donna, con ID tre, non porta né cappello né borsa. L'ultima volta è stata vista attraversare le linee 4 e 3. Non ho altre informazioni su di lei.”

“La persona con ID 8 portano a borsa?”



“No, la persona con ID 8 non porta una borsa.”

“L'utente con ID 10 è maschio o femmina?”

“L'utente con ID dieci è un uomo.”

“Mentre quello con ID 3?”

“L'utente con ID 3 è una donna che non porta né cappello né borsa. È stata vista l'ultima volta attraversare le linee 4 e 3.”

“L'utente con ID 12 ha oltrepassato la linea 1?”

“No, l'utente con ID 12 non ha oltrepassato la linea 1. L'ultima volta è stato visto attraversare la linea 2.”

“Invece quello con ID 8?”

“La persona con ID 8 è un uomo, porta un cappello ma non una borsa. È stato visto attraversare le linee 2, 4, 3, 3, 4.”

“Ho smarrito mio figlio potresti aiutarmi a trovarlo?”

“Certo, mi dispiace per la situazione. Posso aiutarti cercando di individuare tuo figlio tra le persone che ho registrato nel centro commerciale. Per favore, dimmi se tuo figlio porta un cappello o una borsa, e se possibile, quali linee ha attraversato. Più informazioni mi fornisci, meglio potrò aiutarti.”

“Mio figlio ha un cappello”

“Okay, grazie per l'informazione. Ora ho un'informazione in più per aiutarti a ritrovare tuo figlio. Per favore, dimmi anche se porta una borsa e quali linee ha attraversato.”

“Mio figlio non porta una borsa ma non ricordo quale linea è attraversato”

“Capisco la tua preoccupazione. Cercherò di aiutarti a trovare tuo figlio. Considerando che porta un cappello e non una borsa, gli ID che corrispondono a questa descrizione sono: 1, 4, 8, 9, 11, 12, 17 e 18. Per restringere ulteriormente la ricerca, avrei bisogno di sapere quali linee ha attraversato tuo figlio.”

“Guarda ora che me lo stai chiedendo Mi sembra di ricordare che l'ultima volta era vicino alla linea 3”

“Capisco la tua preoccupazione. Considerando che tuo figlio porta un cappello, non una borsa, ed è stato visto vicino alla linea 3, gli ID che corrispondono a questa descrizione sono: 8, 9, 11, 17. Tra questi, l'ID 8 è stato visto attraversare la linea 3 più volte.”

“Va bene Pepper grazie, vado a cercarlo”

“Prego, figurati! Spero che tu riesca a ritrovare tuo figlio presto. Fammi sapere se posso aiutarti ancora.”

“Va bene così grazie, buona giornata”

“Prego, figurati! Sono felice di averti potuto aiutare.”

As we can see, Pepper is able to consistently provide correct answers to the user's questions thanks to the information dynamically retrieved from the file containing data on the people present in the shopping mall. Additionally, it can maintain the context of the conversation by remembering what has been said previously and respond empathetically based on the user's words.

6. Empathy

- Pass: 
- Video:  ./test/Empathy_30L.mp4, ./test/Empathy_Arrabiatissimo.mp4, ./test/Contest_Empathy.mp4, ./test/Peolple_in_the_mall.mp4

“Ciao Pepper

“Ciao! Come posso esserti d'aiuto?”

“Com'è la mia emozione oggi?” [EMOZIONE: Felice]

“Sembri felice!”

“Ho preso 30 e lode a Cognitive Robotics”

“Complimenti per il tuo fantastico risultato! Sono davvero felice per te.”

“Pepper sono arrabbiatissimo”

“Capisco la tua rabbia. Mi dispiace che tu ti senta così. C'è qualcosa che posso fare per aiutarti?”

“...”

“Quale gruppo si è classificato in ultima posizione?”

“Il gruppo quindici si è classificato ultimo, con un F-Score medio (AFS) di 0.488. I suoi membri sono Nappo Domenico, Prugno Siniscalchi Samuel, Valentino Miriam e Parolise Gaetano.”

“Purtroppo dei miei amici facevano parte di quel gruppo”

“Mi dispiace sentirlo. Capisco che possa essere deludente se i tuoi amici non hanno ottenuto il risultato sperato. Posso darti maggiori informazioni sui risultati del contest se lo desideri.”

“...”

“...”

“Ho smarrito mio figlio potresti aiutarmi a trovarlo?”

“Certo, mi dispiace per la situazione. Posso aiutarti cercando di individuare tuo figlio tra le persone che ho registrato nel centro commerciale. Per favore, dimmi se tuo figlio porta un cappello o una borsa, e se possibile, quali linee ha attraversato. Più informazioni mi fornisci, meglio potrò aiutarti.”

“...”



As we can see from the transcription, Pepper can be empathic in two ways: when the question itself conveys information about the emotion or when the emotion is accurately recognized by the emotion recognition module, the dialogue module responds correctly. The only issue is that the emotion recognition module often fails to correctly identify people's emotions, frequently returning a "neutral" result.

7. Other Tests

A large number of additional tests on the LLM can be found in the file `test/TestAssistant.py`.

Test Result Analysis

The following table presents a qualitative estimate of the accuracy of the results obtained during the tests, using a Likert scale 1-5. The accuracy values reflect the system's performance in different scenarios, highlighting its strengths and limitations.

TEST	RESULTS	NOTE	REFERENCE
TRACKING	4.5/5	In some video tests, we can see that Pepper loses track of the subject. However, this only happens when the subject remains silent and other people pass through the scene.	VIDEO: People_in_the_mall.mp4
GENEREAL QUESTION ABOUT PEPPER	5/5		
PEOPLE IN THE MALL	5/5		
CONTEST RESULTS INFO	5/5		
EMPATHY	4/5	When the question itself conveys information about the emotion, the dialogue module responds correctly. The same happens when the emotion is accurately recognized by the emotion recognition module. The only issue is that the emotion recognition module often fails to correctly identify people's emotions, frequently returning a "neutral" result.	
QUESTIONS OUT OF CONTEST	5/5		
CONTEST RETENTION	4.8/5	Pepper is able to maintain the context of previous questions. The only instance of a slight mistake occurred in test 3.	TEST N. 3, responding to question "che punteggio hanno fatto?"

Conclusion

The tests conducted on Pepper demonstrate that it is highly capable of managing conversations in a shopping mall setting. The robot successfully tracks people, retrieves relevant information, maintains conversational context, and even expresses empathy when appropriate.

Pepper excels in answering general questions about itself, providing information about people in the mall, and delivering contest results with high accuracy. Additionally, it correctly handles off-topic questions by explaining its limitations while remaining helpful.

One notable limitation is the emotion recognition module, which frequently misidentifies emotions as "neutral," impacting the system's empathetic responses. Another minor issue is occasional tracking loss in crowded scenes when the subject remains silent.

Overall, Pepper performs exceptionally well in various scenarios, achieving high accuracy scores in almost all test cases. The system effectively integrates information retrieval, context retention, and interactive engagement, making it a valuable assistant in public environments. Further improvements in emotion recognition and tracking consistency could enhance its capabilities even more.