

Conversational AI Agent v0.3 Global Quick Start

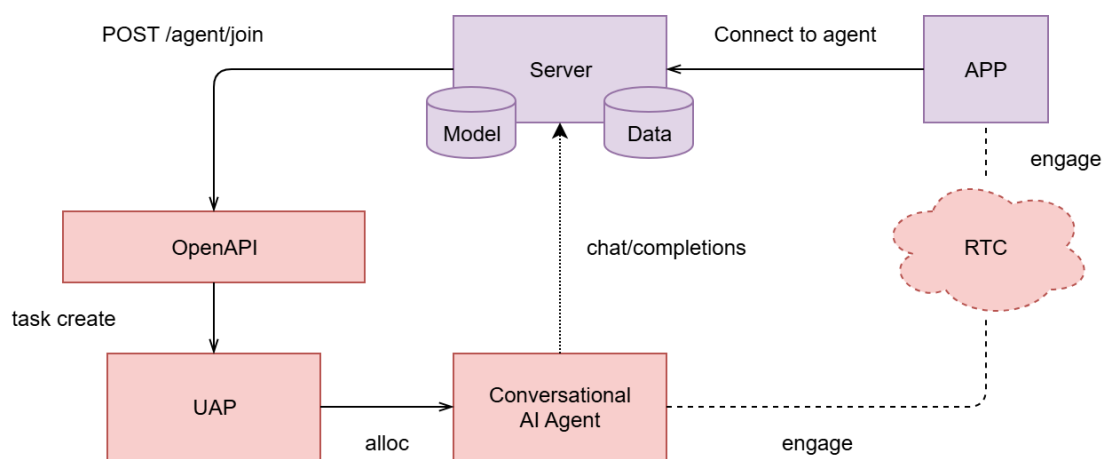
Conversational AI Agent is a next-generation AI Agent solution provided by Agora, designed to help clients simplify real-time interactions with intelligent agents.

Currently, this service is in the Private Beta stage, and the final agreement and API might undergo slight adjustments; your understanding is appreciated.

At present, the service capabilities need to be manually activated. For activation details, please contact [shensiwei](#).

The difference will be marked in **RED** with last version: [Conversational AI Agent v0.2 Global Quick Start](#)

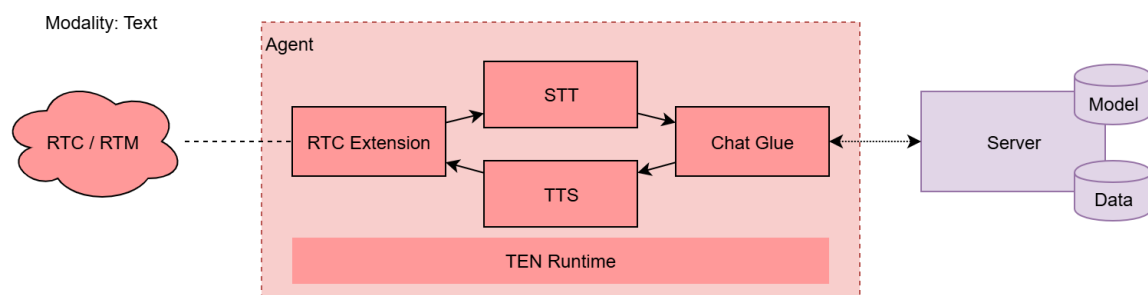
Architecture



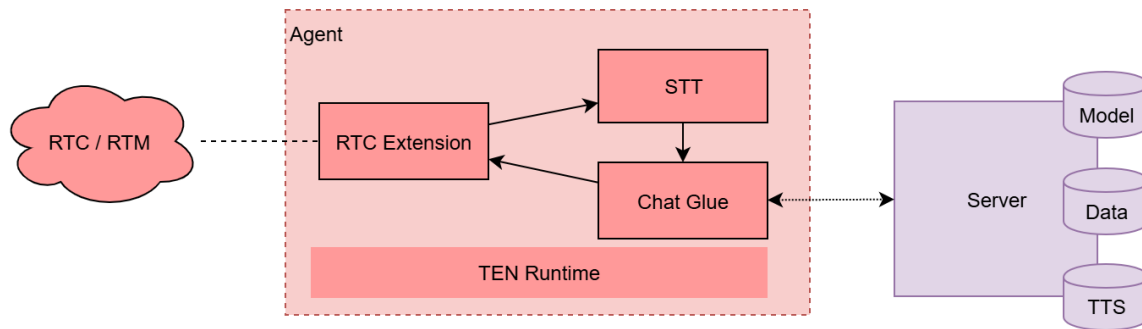
Through the RESTful API provided by Agora, clients can allocate an AI Agent to join a specified RTC channel for 1-on-1 communication with end user.

Note: Does not support multi-user communication.

The Conversational AI Agent includes STT (Speech-to-Text) and TTS (Text-to-Speech) modules. It makes requests to the client's provided URL via an [OpenAI-compatible protocol](#) and retrieves the response in a streaming. The detailed internal implementation is as follows:



For STT and TTS, they are provided by Agora. If clients need to use their own implemented or integrated TTS, they can opt to use the Audio Modal supported by Chat Glue. The specific implementation is as follows:



In this scenario, the input is text, and the output is audio. Chat Completion directly supports the audio protocol used in OpenAI's audio-preview, accepting base64-encoded PCM segments.

API

Agent Join

Method: POST

URL: <https://api.agora.io/api/conversational-ai-agent/v1/projects/{appid}/join>

Parameters:

- Name: Unique identifier for the Agent, identical identifiers cannot be created repeatedly
- Properties: Specific parameters for the Agent, details are as follows

```

{
  "channel": "channel_name",
  "token": "token",
  "agent_rtc_uid": 1234,
  "remote_rtc_uid": 123,
  "input_modalities": ["text"],
  "output_modalities": ["text"],
  "idle_timeout": 120,
  "audio_scenario": 9,
  "rtc_codec": -1,
  "custom_llm": {
    "url": "https://api.openai.com/v1/chat/completions",
    "token": "sk-xxx",
    "prompt": "You are a helpful chatbot.",
    "model": "gpt-4o-mini",
    "greeting": "Merry Christmas, how can I assist you today?",
    "failure_message": "Please hold on a second.",
    "extra_metadata": {
      "userName": "Tomas"
    },
    "metadata_enabled": true
  },
  "tts": {
    "voice_id": "en-US-AndrewMultilingualNeural"
  },
  "asr": {
    "language": "en-US"
  },
  "vad": {
    "silence_duration_ms": 1000
  }
}

```

Descriptions for the parameters:

Property	Type	Description	Default Value	Notes
token	string	rtc channel token	""	
channel	string	rtc channel name	Required	
agent_rtc_uid	int32	rtc int user id	Required	0: random user id [do not forget token] other: specific int user id for agent
remote_rtc_uid	int32	rtc subscription id	Required	0: subscribe all users in channel other: subscribe single user with specific int user id

rtc_codec	int32	agent publish codec	-1	0: PCMU 8: PCMA 9: G722 120: OPUS 122: OPUSFB
audio_scenario	int32	audio scenario	0	0: default 9: AI QOS (only support native sdk with AI QOS)
input_modalities	array[string]	input modalities for custom llm	["text"]	Support ["text"], ["text", "image"]
output_modalities	array[string]	output modalities for custom llm		Support ["audio"], ["text"], ["text", "audio"]
idle_timeout	int	max idle seconds for agent	30	0: agent only stop for manual leave
asr_language	string	asr language	"en-US"	zh-CN, en-US
tts_voice_id	string	tts voice id	"en-US-AndrewMultilingualNeural"	For supported voice IDs, refer to the appendix
custom_llm.url	string	callback url for custom llm	Required	
custom_llm.token	string	verify token for custom llm	""	Must active token in Production Env
custom_llm.prompt	string	prompt for custom llm	""	System prompt
custom_llm.model	string	model id for custom llm	""	
custom_llm.max_history	int32	short term memory length for custom llm	10	0: no short term memory on agent side User and Assistant message will be count seperately
custom_llm.greeting	string	greeting message	""	greet the first rtc user who subscribed for joining the channel
custom_llm.failure_message	string	failure message	""	read out when custom llm call fails
custom_llm.extra_metadata	dict	metadata info	{}	passed to the server via the `metadata` parameter with each custom llm request
custom_llm.metadata_enabled	bool	flag for bypassing metadata	false	<i>Noted that none of the third-party llm service will support this.</i> <i>Enable only when you are hosting your own custom llm service.</i>
vad.interrupt_threshold	float	interruption sensitivity	0.8	0.0 - 1.0
vad.prefix_padding_ms	int	prefix padding in millisecond	300	
vad.silence_duration_ms	int	silence_duration in millisecond	500	
vad.threshold	float	vad sensitivity	0.3	0.0 - 1.0

Success

Status Code: 200

Property	Type	Description
agent_id	string	A unique identifier for the instance, used for subsequent deletion & get.
create_ts	int	create timestamp

state	string	current status
-------	--------	----------------

Failed

StatusCode: non-200

Property	Type	Description
detail	string	Detailed failure info
reason	string	Failure reason

For error codes, refer to the appendix.

Agent Leave

Method: POST

URL: <https://api.agora.io/api/conversational-ai-agent/v1/projects/{appid}/agents/{agentId}/leave>

The agentId should be filled in with the agent_id returned by the join operation.

Agent Update

Allow parameters change for in-progress agent

Method: POST

URL: <https://api.agora.io/cn/api/conversational-ai-agent/v1/projects/{appid}/agents/{agentId}/update>

The agentId should be filled in with the agent_id returned by the join operation, only the following parameters are allow to update:

Property	Type
token	string
tts.voice_id	string
custom_llm.url	string
custom_llm.token	string
custom_llm.prompt	string
custom_llm.model	string
custom_llm.greeting	string
custom_llm.failure_message	string
custom_llm.extra_metadata	dict
custom_llm.metadata_enabled	bool
vad.interrupt_threshold	float
vad.prefix_padding_ms	int
vad.silence_duration_ms	int
vad.threshold	float

Agent Get

Get active agent's status

Method: GET

URL: <https://api.agora.io/cn/api/conversational-ai-agent/v1/projects/{appid}/agents/{agentId}>

Agent List

Fetch agents info that meets the filter

Method: GET

URL: <https://api.agora.io/cn/api/conversational-ai-agent/v1/projects/{appid}/agents>

Glue Sample

Currently, the request requires the client to prepare the Glue URL and Token themselves.

Third-party LLM Service

Currently, Glue's protocol is compatible with OpenAI, so all large model providers compatible with OpenAI can be directly integrated. For example, to access OpenAI Chat Completion directly, you can use the following method:

```
curl --location 'https://api.agora.io/api/conversational-ai-agent/v1/projects/{appid}/join'\
--header 'Authorization: Basic ${key}'\
--header 'Content-Type: application/json'\
--data '{"name":"unique_name","properties":{"channel":"test_channel", "agent_rtc_uid":1234,\
"custom_llm":{"url":"https://api.openai.com/v1/chat/completions", "token":"sk-xxxxx", "model":"gpt-4o-\
mini", "prompt":"You are a helpful assistant.", "greeting":"Hi, how can I assist you today?"}}}'-i
```

DIY

The overall protocol is consistent with [OpenAI](#). Below is a sample code in Python:

```
import os
import json
import traceback
import logging
import asyncio

from typing import List, Union, Dict, Optional
from pydantic import BaseModel, HttpUrl

from fastapi.security import HTTPBearer, HTTPAuthorizationCredentials
from fastapi.responses import StreamingResponse
from fastapi import Depends, FastAPI, HTTPException, Request

logger = logging.getLogger(__name__)

app = FastAPI(title="Chat Completion API",
              description="API for streaming chat completions with support for text, image, and audio content",
              version="1.0.0")

class TextContent(BaseModel):
    type: str = "text"
    text: str

class ImageContent(BaseModel):
    type: str = "image"
    image_url: HttpUrl

class AudioContent(BaseModel):
    type: str = "input_audio"
    input_audio: Dict[str, str]

class ToolFunction(BaseModel):
    name: str
    description: Optional[str]
    parameters: Optional[Dict]
```

```

    strict: bool=False

class Tool(BaseModel):
    type: str="function"
    function: ToolFunction

class ToolChoice(BaseModel):
    type: str="function"
    function: Optional[Dict]

class ResponseFormat(BaseModel):
    type: str="json_schema"
    json_schema: Optional[Dict[str, str]]

class SystemMessage(BaseModel):
    role: str="system"
    content: Union[str, List[str]]

class UserMessage(BaseModel):
    role: str="user"
    content: Union[str, List[Union[TextContent, ImageContent, AudioContent]]]

class AssistantMessage(BaseModel):
    role: str="assistant"
    content: Union[str, List[TextContent]] =None
    audio: Optional[Dict[str, str]] =None
    tool_calls: Optional[List[Dict]] =None

class ToolMessage(BaseModel):
    role: str="tool"
    content: Union[str, List[str]]
    tool_call_id: str

class ChatCompletionRequest(BaseModel):
    context: Optional[Dict] =None
    model: Optional[str] =None
    messages: List[Union[SystemMessage, UserMessage, AssistantMessage, ToolMessage]]
    response_format: Optional[ResponseFormat] =None
    modalities: List[str] =["text"]
    audio: Optional[Dict[str, str]] =None
    tools: Optional[List[Tool]] =None
    tool_choice: Optional[Union[str, ToolChoice]] ="auto"
    parallel_tool_calls: bool=True
    stream: bool=True
    stream_options: Optional[Dict] =None

security =HTTPBearer()

def verify_token(credentials: HTTPAuthorizationCredentials =Depends(security)):
    token =credentials.credentials
    if token !=os.getenv("API_TOKEN"):
        logger.warning("Invalid or missing token")
        raise HTTPException(status_code=403, detail="Invalid or missing token")

@app.post("/chat/completions", dependencies=[Depends(verify_token)])
async def create_chat_completion(request: ChatCompletionRequest, req: Request):
    try:
        logger.debug(f"Received request: {request.json()}")
        #TODO Your logic here
        yield "data: [DONE]\n\n"
    except asyncio.CancelledError:
        logger.info("Request was cancelled")
        raise HTTPException(status_code=499, detail="Request was cancelled")
    except Exception as e:
        traceback_str =''.join(traceback.format_tb(e.__traceback__))
        error_message =f"{str(e)}\n{traceback_str}"
        logger.error(error_message)
        raise HTTPException(status_code=500, detail=error_message)

if __name__ == "__main__":
    import uvicorn
    from fastapi.security import HTTPBearer, HTTPAuthorizationCredentials
    from fastapi import Depends

    uvicorn.run(app, host="0.0.0.0", port=8000)

```

Vision Modality

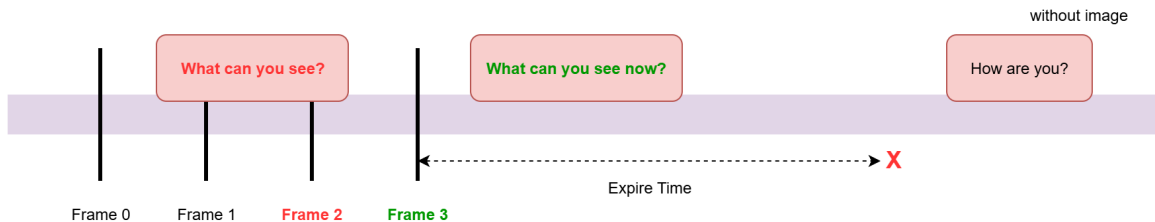
We now support image as a modality, and that need the custom llm to support as well.

The protocol is the same as OpenAI:

```
{
  "role": "user",
  "content": [{
    "type": "image_url",
    "image_url": {
      "url": "jpeg_base64encoded"
    }
  }]
}
```

Agent will record rtc video frame every one second. Each time a new query is to be sent to custom llm, the encoded video frame will be sent along:

- Client can send live video stream, and every query will carry the latest frame received.
- Client can send still image via video stream temporarily when a query of vision modality is need.
- The video frame will expired after 3 seconds.



Transcription

We use datastream to transfer transcription for both user and agent:

- iOS: https://doc.shengwang.cn/api-ref/rtc/ios/API/toc_datastream#callback_irtcengineeventhandler_onstreammessage
- Android: https://doc.shengwang.cn/api-ref/rtc/android/API/toc_datastream#onStreamMessage

It will be sent during the conversation in the same channel:

- is_final: whether the transcription will change later (**Always True**)
- stream_id: rtc int uid for user, 0 for agent and non-0 for user.
- message_id: unique message id
- data_type: "transcribe" as default
- text_ts: timestamp in millisecond
- text: transcription text.

One transcription will be divided into several packets encoded as {message_id} | {id} | {total} | {chunked_data}, and data will be in base64 encoding.

raw

```
{
  "is_final": true,
  "stream_id": 0,
  "message_id": "abcdefgh",
  "data_type": "transcribe",
  "text_ts": 111111111,
  "text": "Once upon a time, in a small village nestled between rolling hills, there was a
mysterious forest known as the Whispering Woods. The villagers often spoke of the woods in hushed
tones, for it was said that the trees could speak, sharing secrets of the past and visions of the
future. However, few dared to enter, fearing the unknown. One day, a curious girl named Elara decided
to explore the woods. She had always been fascinated by the tales told by the village elders. With a
heart full of courage and a sense of adventure, she set off at dawn, leaving behind the familiar
sounds of her village. As Elara stepped into the Whispering Woods, she was greeted by a gentle breeze
that rustled the leaves above. It felt as if the trees were welcoming her. The deeper she ventured,
the more alive the forest seemed. She could hear faint whispers, like a soft melody carried by the
wind. "Who's there?" Elara called out, her voice echoing in the stillness. To her surprise, the
whispers grew louder, forming words that danced around her. "Seek the heart of the woods, and you shall
find what you seek," the trees murmured in unison. Intrigued, Elara followed the sound, her heart
racing with anticipation. After walking for what felt like hours, she came upon a clearing bathed in
sunlight. In the center stood an ancient tree, its trunk wide and gnarled, with branches that reached
towards the sky. It was the oldest tree in the forest, known as the Heartwood. Approaching the
Heartwood, Elara placed her hand on its rough bark. Suddenly, she felt a warmth spreading through her
fingers, and the whispers transformed into a clear voice. "You have shown great courage, young one.
What is it that you seek?" Elara took a deep breath. "I seek knowledge and understanding. I want to
know my purpose in this world." The Heartwood paused, and Elara felt the air around her shimmer. "Your
purpose is not a single path, but a tapestry woven from your choices, your kindness, and your dreams.
Embrace your journey, and you will find your way." With that, the whispers faded, and a sense of peace
enveloped Elara. She realized that her journey was just beginning. Thanking the Heartwood, she made
her way back through the woods, feeling lighter and more connected to the world around her. When she
returned to the village, Elara shared her experience with the villagers. Inspired by her bravery,
others began to explore the Whispering Woods, discovering their own paths and purposes. From that day
on, the village thrived, with each person finding their unique place in the world. And the Whispering
Woods remained a sanctuary of wisdom, welcoming all who sought its secrets. And so, the story of Elara
and the Whispering Woods became a cherished tale, reminding everyone that the journey of self-
discovery is filled with wonder and possibility."
}
```

base64 encoded

ewoJIm1zX2ZpbmFsIjogdHJ1ZSwKCSJzdHJlYW1fawQioiAwLAoJIm1lc3NhZ2VfawQioiAiYWJjZGVmZ2giLaoJImRhIdGFfdHlwZSI6ICJ0cmFuc2NyaWJlIiwKCSJ0ZXh0X3RzIjogMTEeMTExMTExMSwKCSJ0ZXh0IjogIk9uY2UgdXBvbiBhIHRpbWUsIGluIGEGc21hbGwgdm1sbGFnZSBuZWN0bGVkIGJldHdlZW4gcm9sbGluZyBoawXscywgGdGhlcmUgd2FzIGEGbXlzdGVyaW91cyBmb3J1c3Qga25vd24gYXMgdGhlIFdoaxNwZXJpbmcgV29vZHMuIFRoZSB2awxsYwd1cnMgb2Z0ZW4gc3Bva2Ugb2YgdGhlIHdvd2RzIGluIGhlc2hlZCB0b251cywgZm9yIGl0IHdhcyBzYw1kIHRoYXQgdGhlIHRYZWVzIGNvdWxkIHNoZWFrLCBzaGFyaW5nIHNIY3JldHMgb2YgdGhlIHhlc3QgYw5kIHZpc21vbnMgb2YgdGhlIGZ1dHVyZS4gSG93ZXZlc1wgZmV3IGRhcmVkJHRvIGVudGVyLCBmZWfyaW5nIHROZSB1bmtub3duLk9uZSBkYXksIGEGY3Vyaw91cyBnaXJsIG5hbWVkJEVsYXJhIGRlY21kZWQgdG8gZXhwbG9yZSB0aGUgd29vZHMuIFNoZSB0YwQgYwX3YX1zIGJlZW4gZmFzY2luYXRlZCBieSB0aGUgdGFsZXMgdG9sZCBieSB0aGUGdm1sbGFnZSB1bGR1cnMuIFdpcGggYSBoZWfYdCBmdwxsIG9mIGNvdXJhZ2UgYW5kIGEGc2Vuc2Ugb2YgYWR2ZW50dXJlLCBzaGUGc2V0IG9mZiBhdCBkYXduLCBsZWf2aw5nIGJlaGluZCB0aGUGZmFtaWxpYXIGc291bmRzIG9mIGhlc1B2awxsYwd1LkFzIEVsYXJhIHNoZXBwZWQgaw50byB0aGUGV2hpc3B1cm1uZyBxb29kcywgG2h1IHdhcyBncmVldGVkIGJ5IGEGZ2VudGx1IGJyZWV6ZSB0aGF0IHJ1c3RsZWQgdGhlIGx1YXZlc1cyBhYm92ZS4gSXQgZmVsdCBhcyBpZiB0aGUGdHJlZXMgd2VyZSB3ZWxjb21pbmcgaGVyLiBUaGUGZGVlcGVyIHNoZSB2ZW50dXJlZCwgGdGhlIG1vcuUgYwXpdmUgdGhlIGZvcuVzdBzZWVtZWQuIFNoZSBjb3VsZCB0ZWfYIGZhaW50IHdoaxNwZXJzLCBsawt1IGEGc29mdCBtZWxvZHkgY2Fycm1lZCBieSB0aGUGd21uZC7igJxXaG

/igJlZIHROZXL1P+KAnSBfBGfYYSBjYwxsZWQgb3V0LCBoZXIgd9pY2UgZWNoB21uZyBpb1B0aGUGc3RpbGxuZXNzLiBUbyBoZXIGc3VychJpc2UsIHRoZSB3aG1zcGVycyBncmV3IGxvdWRlciwgZm9ybWluZyB3b3JkcyB0aGF0IGRhbmNlZCBhcm91bmQgaGVyLuKAnFNlZWsgdGhlIGh1YXJ0IG9mIHRoZSB3b29kcywgYw5kIHlvdSBzaGFsbCBmaW5kIHdoYXQgew91IHNIZWss4oCdIHRoZSB0cmVlcyBtdXJtdXJlZCBpb1B1bm1zb24uSW50cm1ndWVwKLCBFbGFYYSBmb2xsb3dlZCB0aGUGc291bmQsIGhlc1BoZWfYdCBYyWnNpbmcgd210aCBhbnRpY21wYXRpb24uIEFmdGVyIHdhbGtpbmcgZm9yIHdoYXQgZmVsdCBsaWt1IGhvdXJzLCBzaGUGY2FtZSB1cG9uIGEGY2x1YXJpbmcgYmF0aGwKIGluIHNI1bmXpZ2h0LiBjb1B0aGUGY2VudGVyIHNoB29kIGFuIGFuY211bnQgdHJlZSwgaXRzIHRydW5rIHdpZGUGYw5kIGduYXJ5SjZwQsIHdpdGggYnJhbmNoZXMgdGhhdCBYzWfjaGVkIHRvd2FyZHMgdGhlIHnreS4gSXQgd2FzIHROZSBvbGRlc3QgdHJlZSBpb1B0aGUGZm9yZXN0LCBrbm93biBhcyB0aGUGSGVhcnR3b29kLkFwcHJvYWN0aW5nIHRoZSBIZWfYdHdvb2QsIEVsYXJhIHBSYWNlZCB0ZXIgaGFuZCBvb1BpdHMgcm91Z2ggYmFyay4gU3VkcGVubHksIHNoZSBmZWx0IGEGd2FybXRoIHNoZSBmVhZGluZyB0aHJvdWdoIGhlc1BmaW5nZXJzLCBhbmQgdGhlIHdoaxNwZXJzIHRyYw5zZm9ybWVkJGludG8gYSBjbGVhcyB2b21jZS4g4oCcWw91IGhhdMUGc2hvd24gZ3JlYXQgY291cmFnZSwgew91bmcbg251LiBxaGF0IGlZlIGl0IHROYXQgew91IHNIZWss

/4oCdRwxhcmEgdG9vayBhIGRlZXAgYnJlYXRoLiDigJxJIHNlZWsga25vd2x1ZGdlIGFuZCB1bmR1cnN0Yw5kaw5nLiBjIHdhbnQgdG8ga25vdyBteSBwdXJwb3NlIGluIHRoaXMgd29ybGQu4oCdVGHlIEh1YXJ0d29vZCBwYXVzZWQsIGFuZCBFbGFYYSBmZWx0IHRoZSBhaXIgYXJvdW5kIGhlc1BzaGltbWVYLiDigJxZb3VyIHb1cnBvc2UgaXMGbm90IGEGc2luZ2x1IHbhdGgsIGJ1dCBhIHRhcGVzdHJ5IHdvdmVuIGZyb20gew91ciBjaG9pY2VzLCB5b3VyIGtpbmRuZXNzLCBhbmQgew91ciBkcmVhbXMuIEVtYnJhY2Ugew91ciBqb3VybmV5LCBhbmQgew91IHdpbGwgZmluZCB5b3VyIHdheS7igJ1XaXRoIHRoYXQsIHRoZSB3aG1zcGVycyBmYWRlZCwgYw5kIGEGc2Vuc2Ugb2YgcGvhY2UgZW52ZWxvcGVkIEVsYXJhLiBtAGUGcmVhbG16ZWQgdGhhdCB0ZXIga91cm51eSB3YXMGanVzdCBiZWdpbm5pbmcuIFRoYw5raW5nIHRoZSBIZWfYdHdvb2QsIHNoZSBtYWRlIGhlc1B3YXkgYmFjayB0aHJvdWdoIHRoZSB3b29kcywgZmVlbGluZyBsaWdodGVyIGFuZCBtb3JlIGNvbm51Y3RlZCB0byB0aGUGd29ybGQgYXJvdW5kIGhlc15XaGVuIHNoZSBYXR1cm51ZCB0byB0aGUGdm1sbGFnZSwgRwxhcmEgc2hhcmVkJGhlc1BleHB1cmllbmNlIHdpdGggdGhlIHZpbGxhZ2Vycy4gSW5zcGlyZWQgYnkgYGVyIGJyYXZlcnkSI90aGVycyBiZWdhbiB0byBleHBsb3JlIHRoZSBXaG1zcGVyaW5nIFdvd2RzLCBkaXNjb3ZlcmluZyB0aGVpc1Bvd24gcGF0aHMgYw5kIHb1cnBvc2VzLkZyb20gdGhhdCBkYXkgb24sIHRoZSB2awxsYwd1IHRocml2ZWQsIHdpdGggZWfjaCBwZXJzb24gZmluZGluZyB0aGVpc1B1bm1xdWUGcGxhY2Ugaw4gdGhlIHdvcmxkLiBBbmQgdGhlIFdoaxNwZXJpbmcgV29vZHMgcmtVYwluZWQgYSBzYw5jdHVhcnkgb2Ygd21zZG9tLCB3ZWxjb21pbmcgYwxsIHdobyBzb3VnaHQgaXRzIHNIY3JldHMuQW5kIHNVLCB0aGUGc3Rvcnkqb2YgRWxhcmEgYw5kIHRoZSBXaG1zcGVyaW5nIFdvd2RzIGJlY2FtZSBhIGNoZXJpc2hlZCB0YwX1LCByZW1pbmRpbmcgZXZlcnlvbmUgdGhhdCB0aGUGam91cm51eSBvZlBzZWxmLWRpc2NvdmVyeSBpcyBmaWxsZWQgd210aCB3b25kZXIgyw5kIHbvc3NpYm1saXR5LiIKfQ==

after divided in 5 packets

abcdefgh|0|5|ewoJIm1zX2Zpb...
abcdefgh|1|5|ZWF2aw5nI...
abcdefgh|2|5|BoZWfYdCBYw...
abcdefgh|3|5|G8ga25vdyBt...
abcdefgh|4|5|d24gcGF0aHMgY...

Appendix

Voice ID

- en-US-AvaMultilingualNeural
- en-US-AndrewMultilingualNeural
- en-US-EmmaMultilingualNeural
- en-US-BrianMultilingualNeural
- en-US-AvaNeural
- en-US-AndrewNeural
- en-US-EmmaNeural
- en-US-BrianNeural
- en-US-JennyNeural
- en-US-GuyNeural
- en-US-AriaNeural
- en-US-DavisNeural
- en-US-JaneNeural
- en-US-JasonNeural
- en-US-SaraNeural
- en-US-TonyNeural
- en-US-NancyNeural

- en-US-AmberNeural
- en-US-AnaNeural
- en-US-AshleyNeural
- en-US-BrandonNeural
- en-US-ChristopherNeural
- en-US-CoraNeural
- en-US-ElizabethNeural
- en-US-EricNeural
- en-US-JacobNeural
- en-US-JennyMultilingualNeural4
- en-US-MichelleNeural
- en-US-MonicaNeural
- en-US-RogerNeural

ErrorCode

Code	Description
400	Invalid request parameters
403	Unauthorized access
409	Agent conflict
422	Access limit exceeded
502	Gateway error
503	Agent startup error
504	Timeout error