

(12) 按照专利合作条约所公布的国际申请

(19) 世界知识产权组织
国际局

(43) 国际公布日
2021 年 3 月 4 日 (04.03.2021)



(10) 国际公布号
WO 2021/036644 A1

(51) 国际专利分类号:
G10L 15/02 (2006.01) *G10L 15/25* (2013.01)
G10L 15/22 (2006.01)

(21) 国际申请号: PCT/CN2020/105046

(22) 国际申请日: 2020 年 7 月 28 日 (28.07.2020)

(25) 申请语言: 中文

(26) 公布语言: 中文

(30) 优先权:
201910820742.1 2019年8月29日 (29.08.2019) CN

(71) 申请人: 腾讯科技(深圳)有限公司 (TENCENT TECHNOLOGY (SHENZHEN) COMPANY LIMITED) [CN/CN]; 中国广东省深圳市南山区高新区科技中一路腾讯大厦 35 层, Guangdong 518057 (CN)。

(72) 发明人: 康世胤(KANG, Shiyin); 中国广东省深圳市南山区高新区科技中一路腾讯大厦 35 层, Guangdong 518057 (CN)。 陀得意(TUO, Deyi); 中国广东省深圳市南山区高新区科技中一路腾讯大厦 35 层, Guangdong 518057 (CN)。 李广之(LEI, Kuongchi); 中国广东省深圳市南山区高新区科技中一路腾讯大厦 35 层, Guangdong 518057 (CN)。 傅天晓(FU, Tianxiao); 中国广东省深圳市南山区高新区科技中一路腾讯大厦 35 层, Guangdong 518057 (CN)。 黄晖榕(HUANG, Huirong); 中国广东省深圳市南山区高新区科技中一路腾讯大厦 35 层, Guangdong 518057 (CN)。 苏丹(SU, Dan); 中国广东省深圳市南山区高新区科技中一路腾讯大厦 35 层, Guangdong 518057 (CN)。

(74) 代理人: 深圳市深佳知识产权代理事务所(普通合伙) (SHENPAT INTELLECTUAL PROPERTY

(54) Title: VOICE-DRIVEN ANIMATION METHOD AND APPARATUS BASED ON ARTIFICIAL INTELLIGENCE

(54) 发明名称: 一种基于人工智能的语音驱动动画方法和装置

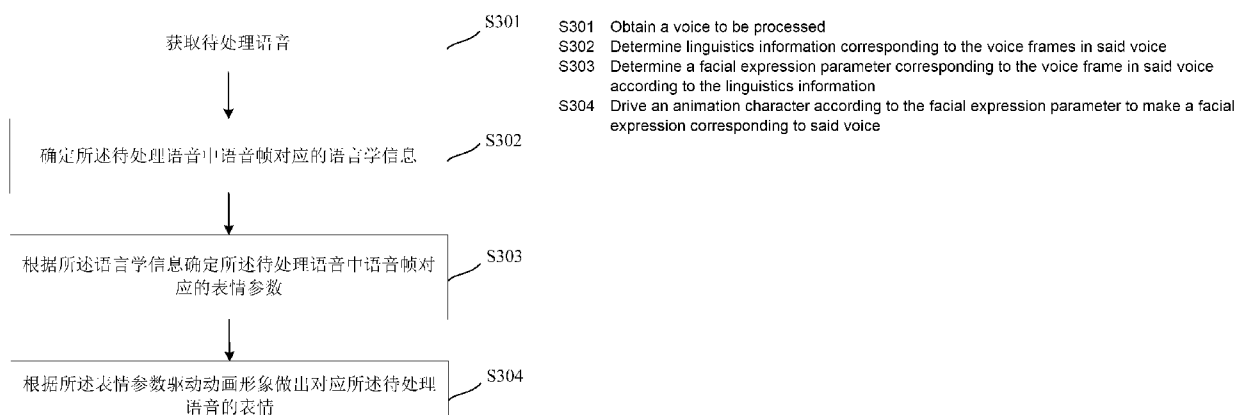


图 3

(57) Abstract: A voice-driven animation method and apparatus based on artificial intelligence. The method comprises: obtaining a voice to be processed (S301), said voice comprising a plurality of voice frames; determining linguistics information corresponding to the voice frames in said voice (S302), the linguistic information being used for identifying a probability of distribution of phonemes to which the corresponding voice frame in said voice belongs; determining a facial expression parameter corresponding to the voice frame in said voice according to the linguistics information (S303); and driving an animation character according to the facial expression parameter to make a facial expression corresponding to said voice (S304).

(57) 摘要: 一种基于人工智能的语音驱动动画方法和装置, 该方法包括: 获取待处理语音 (S301), 待处理语音包括多个语音帧; 确定待处理语音中语音帧对应的语言学信息 (S302), 语言学信息用于标识待处理语音中语音帧所属音素的分布可能性; 根据语言学信息确定待处理语音中语音帧对应的表情参数 (S303); 根据表情参数驱动动画形象做出对应待处理语音的表情 (S304)。

AGENCY): 中国广东省深圳市罗湖区南湖街道春风路庐山大厦B座18C2、18D、18E、18E2, Guangdong 518001 (CN)。

(81) 指定国 (除另有指明, 要求每一种可提供的国家保护): AE, AG, AL, AM, AO, AT, AU, AZ, BA, BB, BG, BH, BN, BR, BW, BY, BZ, CA, CH, CL, CN, CO, CR, CU, CZ, DE, DJ, DK, DM, DO, DZ, EC, EE, EG, ES, FI, GB, GD, GE, GH, GM, GT, HN, HR, HU, ID, IL, IN, IR, IS, IT, JO, JP, KE, KG, KH, KN, KP, KR, KW, KZ, LA, LC, LK, LR, LS, LU, LY, MA, MD, ME, MG, MK, MN, MW, MX, MY, MZ, NA, NG, NI, NO, NZ, OM, PA, PE, PG, PH, PL, PT, QA, RO, RS, RU, RW, SA, SC, SD, SE, SG, SK, SL, ST, SV, SY, TH, TJ, TM, TN, TR, TT, TZ, UA, UG, US, UZ, VC, VN, WS, ZA, ZM, ZW。

(84) 指定国 (除另有指明, 要求每一种可提供的地区保护): ARIPO (BW, GH, GM, KE, LR, LS, MW, MZ, NA, RW, SD, SL, ST, SZ, TZ, UG, ZM, ZW), 欧亚 (AM, AZ, BY, KG, KZ, RU, TJ, TM), 欧洲 (AL, AT, BE, BG, CH, CY, CZ, DE, DK, EE, ES, FI, FR, GB, GR, HR, HU, IE, IS, IT, LT, LU, LV, MC, MK, MT, NL, NO, PL, PT, RO, RS, SE, SI, SK, SM, TR), OAPI (BF, BJ, CF, CG, CI, CM, GA, GN, GQ, GW, KM, ML, MR, NE, SN, TD, TG)。

本国际公布:

— 包括国际检索报告 (条约第21条(3))。

一种基于人工智能的语音驱动动画方法和装置

本申请要求于 2019 年 8 月 29 日提交中国专利局、申请号 201910820742.1、申请名称为“一种基于人工智能的语音驱动动画方法和装置”的中国专利申请的优先权，其全部内容通过引用结合在本申请中。

5 技术领域

本申请涉及数据处理领域，特别是涉及基于人工智能的语音驱动动画。

背景技术

10 目前，语音到虚拟人脸动画的生成这一技术正在成为工业界应用领域的研究热点，例如针对一段任意说话人的语音，可以驱动一个动画形象做出该段语音对应的口型。在这一场景下，动画形象的存在能极大地增强真实感，提升表现力，带给用户更加沉浸式的体验。

一种方式是通过 Speech2Face 系统实现上述技术。一般来说，针对一说话人的语音，该系统提取语音中的声学特征例如梅尔频率倒谱系数（Mel Frequency Cepstral Coefficient, MFCC）后，通过映射模型可以基于声学特征确定出一个能够调整的动画形象的表情参数，依据该表情参数可以控制该动画形象做出该段语音对应的口型。

15 然而，由于提取的声学特征中含有与说话人相关的信息，导致以此建立的映射模型对特定说话人的语音可以准确确定对应的表情参数，若说话人出现了变更，映射模型确定出的表情参数将出现较大偏差，以此驱动的动画形象口型将与语音不一致，降低了交互体验。

发明内容

20 为了解决上述技术问题，本申请提供了基于人工智能的语音驱动动画方法和装置，可以有效支持任意说话人对应的待处理语音，提高了交互体验。

本申请实施例公开了如下技术方案：

第一方面，本申请实施例提供一种语音驱动动画方法，所述方法由音视频处理设备执行，所述方法包括：

获取待处理语音，所述待处理语音包括多个语音帧；

25 确定所述待处理语音中语音帧对应的语言学信息，所述语言学信息用于标识所述待处理语音中语音帧所属音素的分布可能性；

根据所述语言学信息确定所述待处理语音中语音帧对应的表情参数；

根据所述表情参数驱动动画形象做出对应所述待处理语音的表情。

—2—

第二方面，本申请实施例提供一种语音驱动动画装置，所述装置部署在音视频处理设备
5 上，所述装置包括获取单元、第一确定单元、第二确定单元和驱动单元：

所述获取单元，用于获取待处理语音，所述待处理语音包括多个语音帧；

所述第一确定单元，用于确定所述待处理语音中语音帧对应的语言学信息，所述语言
5 学信息用于标识所述待处理语音中语音帧所属音素的分布可能性；

所述第二确定单元，用于根据所述语言学信息确定所述待处理语音中语音帧对应的表
情参数；

所述驱动单元，用于根据所述表情参数驱动动画形象做出对应所述待处理语音的表情。

第三方面，本申请实施例提供一种用于语音驱动动画的设备，所述设备包括处理器以
10 及存储器：

所述存储器用于存储程序代码，并将所述程序代码传输给所述处理器；

所述处理器用于根据所述程序代码中的指令执行第一方面所述的方法。

第四方面，本申请实施例提供一种计算机可读存储介质，所述计算机可读存储介质用
于存储程序代码，所述程序代码用于执行第一方面所述的方法。

15 由上述技术方案可以看出，当获取到包括多个语音帧的待处理语音时，可以确定出待
处理语音中语音帧对应的语言学信息，每一个语言学信息用于标识所对应语音帧所属音素
的分布可能性，即体现语音帧中内容属于哪一种音素的概率分布，该语言学信息携带的信
息与待处理语音的实际说话人无关，由此可以抵消不同说话人发音习惯对后续表情参数
20 的确定所带来的影响，根据语言学信息所确定出的表情参数，可以准确驱动动画形象做出对
应待处理语音的表情，例如口型，从而可以有效支持任意说话人对应的待处理语音，提高
了交互体验。

附图说明

25 为了更清楚地说明本申请实施例或现有技术中的技术方案，下面将对实施例或现有技术
描述中所需要使用的附图作简单地介绍，显而易见地，下面描述中的附图仅仅是本申请
的一些实施例，对于本领域普通技术人员来讲，在不付出创造性劳动性的前提下，还可以
根据这些附图获得其他的附图。

图 1 为相关技术中所采用的 Speech2Face 系统；

图 2 为本申请实施例提供了一种语音驱动动画方法的应用场景示例图；

图 3 为本申请实施例提供了一种语音驱动动画方法的流程图；

30 图 4 为本申请实施例提供的 Speech2Face 的系统架构示例图；

图 5 为本申请实施例提供的 ASR 模型的训练过程示例图；

图 6 为本申请实施例提供的基于 DNN 模型确定表情参数的示例图；

图 7 为本申请实施例提供的相邻上下文语音帧的示例图；

图 8 为本申请实施例提供的间隔上下文语音帧的示例图；

图 9 为本申请实施例提供的基于 LSTM 模型确定表情参数的示例图；

5 图 10 为本申请实施例提供的基于 BLSTM 模型确定表情参数的示例图；

图 11 为本申请实施例提供的一种语音驱动动画装置的结构图；

图 12 为本申请实施例提供的终端设备的结构图；

图 13 为本申请实施例提供的服务器的结构图。

具体实施方式

10 下面结合附图，对本申请的实施例进行描述。

相关技术中所采用的 Speech2Face 系统参见图 1 所示。针对一说话人的语音，该系统可以对语音进行声学特征提取，得到 MFCC。然后，通过映射模型基于声学特征确定出表情参数。对于设定好的、可以通过调整表情参数来调整表情（例如口型）的动画形象，利用确定出的表情参数调整动画形象，生成该段语音对应的动画形象。

15 然而，由于相关技术中提取的声学特征与说话人相关，当说话人更换时，映射模型确定出的表情参数将出现较大偏差，以此驱动的动画形象的口型将与语音不一致，降低了交互体验。

为此，本申请实施例提供一种基于人工智能的语音驱动动画方法，该方法在获取到包括多个语音帧的待处理语音后，可以确定出待处理语音中语音帧对应的语言学信息。与相关技术中提取的声学特征例如 MFCC 相比，语言学信息携带的信息与待处理语音对应的实际说话人无关，避免了不同说话人发音习惯对后续表情参数的确定所带来的影响，故，针对任意说话人对应的待处理语音，可以根据语言学信息确定表情参数，从而准确驱动动画形象做出对应待处理语音的表情。

20 需要强调的是，本申请实施例所提供的语音驱动动画方法是基于人工智能实现的，人工智能（Artificial Intelligence, AI）是利用数字计算机或者数字计算机控制的机器模拟、延伸和扩展人的智能，感知环境、获取知识并使用知识获得最佳结果的理论、方法、技术及应用系统。换句话说，人工智能是计算机科学的一个综合技术，它企图了解智能的实质，并生产出一种新的能以人类智能相似的方式做出反应的智能机器。人工智能也就是研究各种智能机器的设计原理与实现方法，使机器具有感知、推理与决策的功能。

25 人工智能技术是一门综合学科，涉及领域广泛，既有硬件层面的技术也有软件层面的技术。人工智能基础技术一般包括如传感器、专用人工智能芯片、云计算、分布式存储、

—4—

大数据处理技术、操作/交互系统、机电一体化等技术。人工智能软件技术主要包括计算机视觉技术、语音处理技术、自然语言处理技术以及机器学习/深度学习等几大方向。

在本申请实施例中，主要涉及的人工智能软件技术包括上述语音处理技术和机器学习等方向。

5 例如可以涉及语音技术(Speech Technology)中的语音识别技术，其中包括语音信号预处理(Speech signal preprocessing)、语音信号频域分析(Speech signal frequency analyzing)、语音信号特征提取(Speech signal feature extraction)、语音信号特征匹配/识别(Speech signal feature matching / recognition)、语音的训练(Speech training)等。

10 例如可以涉及机器学习(Machine learning, ML)，机器学习是一门多领域交叉学科，涉及概率论、统计学、逼近论、凸分析、算法复杂度理论等多门学科。专门研究计算机怎样模拟或实现人类的学习行为，以获取新的知识或技能，重新组织已有的知识结构使之不断改善自身的性能。机器学习是人工智能的核心，是使计算机具有智能的根本途径，其应用遍及人工智能的各个领域。机器学习通常包括深度学习(Deep Learning)等技术，深度学习包括人工神经网络(artificial neural network)，例如卷积神经网络(Convolutional Neural
15 Network, CNN)、循环神经网络(Recurrent Neural Network, RNN)、深度神经网络(Deep neural network, DNN)等。

本申请实施例提供的基于人工智能的语音驱动动画方法可以应用于具有驱动动画能力的音视频处理设备，该音视频处理设备可以是终端设备，也可以是服务器。

20 该音视频处理设备可以具有实施语音技术中自动语音识别技术(ASR)和声纹识别的能力。让音视频处理设备能听、能看、能感觉，是未来人机交互的发展方向，其中语音成为未来最被看好的人机交互方式之一。

在本申请实施例中，音视频处理设备通过实施上述语音技术，可以对获取的待处理语音进行识别，以及确定待处理语音中语音帧对应的语言学信息等功能；通过机器学习技术训练神经网络映射模型，通过训练得到的神经网络映射模型根据语言学信息确定表情参数，
25 以驱动动画形象做出对应待处理语音的表情。

其中，若音视频处理设备是终端设备，则终端设备可以是智能终端、计算机、个人数字助理(Personal Digital Assistant, 简称PDA)、平板电脑等。

若该音视频处理设备是服务器，则服务器可以为独立服务器，也可以为集群服务器。当服务器实施该语音驱动动画方法时，服务器确定出表情参数，利用该表情参数驱动终端
30 设备上的动画形象做出对应待处理语音的表情。

需要说明的是，本申请实施例提供的语音驱动动画方法可以应用到多种承担一些人类工作的应用场景，例如新闻播报、天气预报以及游戏解说等，还能承担一些私人化的服务，例如心理医生，虚拟助手等面向个人的一对一服务。在这些场景下，利用本申请实施例提供的方法驱动动画形象做出表情，极大地增强真实感，提升表现力。

为了便于理解本申请的技术方案，下面结合实际应用场景对本申请实施例提供的语音驱动动画方法进行介绍。

参见图 2，图 2 为本申请实施例提供的基于人工智能的语音驱动动画方法的应用场景示例图。该应用场景以音视频处理设备为终端设备为例进行介绍。该应用场景中包括终端设备 201，终端设备 201 可以获取待处理语音，待处理语音中包括多个语音帧。待处理语音为与任意说话人对应的语音，例如为说话人发出的语音。本实施例并不限定待处理语音的类型，其中，待处理语音可以是说话人说话所对应的语音，也可以是说话人唱歌所对应的语音。本实施例也不限定待处理语音的语种，例如待处理语音可以是汉语、英语等。

可以理解的是，待处理语音不仅可以为说话人通过终端设备 201 输入的语音。在一些情况下，本申请实施例提供的方法所针对的待处理语音也可以是根据文本生成的语音，即通过终端设备 201 输入文本，通过智能语音平台转换成符合说话人语音特点的语音，将该语音作为待处理语音。

由于音素是根据语音的自然属性划分出来的最小语音单位，依据音节里的发音动作来分析，一个动作（例如口型）构成一个音素。也就是说，音素与说话人无关，无论说话人是谁、无论待处理语音是英语还是汉语、无论发出音素所对应的文本是否相同，只要待处理语音中语音帧对应的音素相同，那么，对应的表情例如口型具有一致性。基于音素的特点，在本实施例中，终端设备 201 可以确定待处理语音中语音帧对应的语言学信息，语言学信息用于标识待处理语音中语音帧所属音素的分布可能性，即语音帧内容属于哪一种音素的概率分布，从而确定语音帧内容所属音素。

可见，与前述相关内容中所涉及的相关技术中的声学特征相比，语言学信息携带的信息与待处理语音的实际说话人无关，不管说话人是谁、语音帧内容所对应的文本是什么，语音帧内容所属的音素（语言学信息）是可以确定的，例如确定出语音帧内容所属的音素是“a”，虽然音素“a”对应的说话人可能不同，对应的文本也有可能不同，但是，只要发出的是音素“a”，音素“a”对应的表情例如口型是一致的。故，终端设备 201 可以根据语言学信息准确地确定出表情参数，从而准确驱动动画形象做出对应待处理语音的表情，避免了不同说话人发音习惯对表情参数的确定所带来的影响，提高了交互体验。

接下来，将结合附图对本申请实施例提供的语音驱动动画方法进行详细介绍。

参见图 3，图 3 示出了一种语音驱动动画方法的流程图，所述方法包括：

S301、获取待处理语音。

以音视频处理设备是终端设备为例，当说话人通过麦克风向终端设备输入待处理语音，以希望根据待处理语音驱动动画形象做出对应待处理语音的表情时，终端设备可以获取该待处理语音。本申请实施例对说话人、待处理语音类型、语种等不做限定。该方法可以支持任意说话人对应的语音，即任意说话人对应的语音可以作为待处理语音；该方法可以支

持多种语种，即待处理语音的语种可以是汉语、英语、法语等各种语种；该方法还可以支持唱歌语音，即待处理语音可以是说话人唱歌的语音。

S302、确定所述待处理语音中语音帧对应的语言学信息。

5 终端设备可以对待处理语音进行语言学信息提取，从而确定出待处理语音中语音帧对应的语言学信息。其中，语言学信息是与说话人无关的信息，用于标识所述待处理语音中语音帧所属音素的分布可能性。在本实施例中，语言学信息可以包括音素后验概率(Phonetic Posterior grams, PPG)、瓶颈(bottomneck)特征和嵌入(imbedding)特征中任意一种或多种的组合。后续实施例中主要以语言学信息是 PPG 为例进行介绍。

10 需要说明的是，本申请实施例实现语音驱动动画时，所采用的系统也是 Speech2Face 系统，但是本申请实施例采用的 Speech2Face 系统与前述内容相关技术中的 Speech2Face 系统有所不同，本申请实施例提供的 Speech2Face 系统的系统架构可以参考图 4 所示。其中，该 Speech2Face 系统主要由四部分组成，第一部分是训练得到自动语音识别(Automatic Speech Recognition, ASR)模型，用于 PPG 提取；第二部分是基于训练好的 ASR 模型，确定待处理语音的 PPG(例如 S302)；第三部分则是声学参数到表情参数的映射，即基于
15 音素后验概率 PPG 确定待处理语音中语音帧对应的表情参数(例如 S303)，图 4 示出了通过神经网络映射模型，根据 PPG 确定待处理语音中语音帧对应的表情参数的示例；第四部分是根据表情参数驱动设计好的 3D 动画形象做出对应待处理语音的表情(例如 S304)。

20 本实施例中所涉及的音素可以包括 218 种，包括汉语音素、英语音素等多个语种，从而实现多语种的语言学信息提取。若语言学信息是 PPG，得到的 PPG 可以是一个 218 维的向量。

在一种实现方式中，终端设备可以通过 ASR 模型实现语言学信息提取。在这种情况下，以语言学信息是 PPG 为例，为了能够提取 PPG，需要预先训练一个 ASR 模型(即上述第一部分)。ASR 模型可以是根据包括了语音片段和音素对应关系的训练样本训练得到的。在实际训练过程中，ASR 模型是基于 Kaldi 所提供的 ASR 接口，在给定 ASR 数据集的情况下进行训练得到，ASR 数据集中包括了训练样本。其中，Kaldi 是一种开源的语音识别工具箱，Kaldi 使用一个基于深度置信网络-深度神经网络(Deep Belief Network-Deep neural network, DBN-DNN)的网络结构来根据提取的 MFCC 预测语音帧属于每个音素的可能性大小，也即 PPG，从而对输出的音素进行分类，ASR 模型的训练过程可以参见图 5 中虚线所示。当训练得到上述 ASR 模型之后，对于待处理语音，经过上述 ASR 模型之后，就能
25 够输出得到待处理语音中语音帧对应的 PPG，从而可以用于后续表情参数的确定。其中，ASR 模型实际上可以是 DNN 模型。

基于 ASR 模型的训练方式，利用 ASR 模型确定语言学信息的方式可以是确定待处理语音中语音帧对应的声学特征，该声学特征为前述相关内容中所涉及的相关技术中的声学特征，例如 MFCC。然后利用 ASR 模型确定该声学特征对应的语言学信息。

—7—

需要说明的是，由于训练 ASR 模型所使用的 ASR 数据集就考虑了带噪声的语音片段情况，对噪声的适应性更强，因此通过 ASR 模型提取的语言学信息相比 MFCC 等相关技术中所使用的声学特征来说鲁棒性更强。

S303、根据所述语言学信息确定所述待处理语音中语音帧对应的表情参数。

5 表情参数用于驱动动画形象做出对应待处理语音的表情，即通过确定出的表情参数对预先建立的动画形象的表情参数进行调整，从而使得动画形象做出与说出待处理语音相符的表情。

10 通常情况下，对于动画形象而言，表情可以包括面部表情和身体姿态表情，面部表情例如可以包括口型、五官动作和头部姿态等，身体姿态表情例如可以包括身体动作、手势动作以及走路姿态等等。

在一种实现方式中，S303 的实现方式可以是根据语言学信息，通过神经网络映射模型确定待处理语音中语音帧对应的表情参数，例如图 4 中虚线框所示。其中，神经网络映射模型可以包括 DNN 模型、长短期记忆网络（Long Short-Term Memory, LSTM）模型或双向长短期记忆网络(Bidirectional Long Short-term Memory, BLSTM)模型。

15 神经网络映射模型是预先训练得到的，神经网络映射模型实现的是语言学信息到表情参数的映射，即当输入语言学信息时，便可以输出待处理语音中语音帧对应的表情参数。

20 由于语言学信息与表情中的口型较为相关，通过语言学信息确定口型更为精确。而口型之外的其他表情与待处理语音对应的情感更为相关，为了准确的确定表情参数，以便通过表情参数驱动动画形象做出更为丰富的表情，例如，动画形象在做出口型的同时还可以大笑、眨眼等。因此，S303 的一种可能实现方式可以参见图 5 所示（图 5 以语言学信息是 PPG 为例），在利用训练好的 ASR 模型得到语音帧的 PPG 之后，将 PPG（即语言学信息）与预先标注好的情感向量进行拼接得到最终的特征，从而根据 PPG 和待处理语音对应的情感向量，确定待处理语音中语音帧对应的表情参数。

25 其中，在情感方面，本实施例中采用了四种常用的情感，包括快乐、悲伤、生气和正常状态。利用情感向量来表征情感，情感向量采用了 1-of-K 编码方式，即设定长度为 4，在四个维度上分别取 1，其他维度取 0 得到四个向量，用以分别表示四种情感。在确定出语音帧的 PPG 后，即得到一个 218 维的向量，与该待处理语音的 4 维情感向量进行拼接就能得到一个 222 维的特征向量，用于后续作为神经网络映射模型的输入。

30 需要说明的是，在本实施例中所使用的基于神经网络的神经网络映射模型，还可以使用 Tacotron decoder 进行替换。其中，Tacotron decoder 是一种注意力模型，用于端到端的语音合成。

S304、根据所述表情参数驱动动画形象做出对应所述待处理语音的表情。

对于动画形象本身而言，动画形象可以是 3D 的形象也可以是 2D 的形象，本实施例对此不做限定。例如，建立的动画形象如图 2 中的动画形象所示，利用该动画形象向大家拜年，假设当待处理语音为“恭喜发财，大吉大利”，那么，根据该待处理语音可以确定出表情参数，从而驱动该动画形象做出可以发出“恭喜发财，大吉大利”的表情（口型）。

5 当然，还可以做出除了口型之外的其他表情。

由上述技术方案可以看出，当获取到包括多个语音帧的待处理语音时，可以确定出待处理语音中语音帧对应的语言学信息，每一个语言学信息用于标识所对应语音帧所属音素的分布可能性，即体现语音帧中内容属于哪一种音素的概率分布，该语言学信息携带的信息与待处理语音的实际说话人无关，由此可以抵消不同说话人发音习惯对后续表情参数的确定所带来的影响，根据语言学信息所确定出的表情参数，可以准确驱动动画形象做出对应待处理语音的表情，例如口型，从而可以有效支持任意说话人对应的待处理语音，提高了交互体验。

10 在一种可能的实现方式中，待处理语音中包括多个语音帧，接下来将以待处理语音中的一个语音帧作为目标语音帧为例，详细介绍针对目标语音帧，S303 如何确定出目标语音帧对应的表情参数。

由于语音存在协同发音的效应，目标语音帧对应的口型等表情会与其上下文语音帧短时相关，因此，为了可以准确的确定出目标语音帧对应的表情参数，可以在确定表情参数时结合上下文语音帧，即确定目标语音帧所处的语音帧集合，该语音帧集合中包括了目标语音帧和目标语音帧的上下文语音帧，从而根据语音帧集合中语音帧分别对应的语言学信息确定目标语音帧对应的表情参数。

可以理解的是，若使用神经网络映射模型确定表情参数，根据所使用的神经网络映射模型的不同，确定语音帧集合的方式以及语音帧集合中语音帧的数量有所不同。若神经网络映射模型是 DNN 模型，即基于前向连接的分类器，如图 6 所示，由于 DNN 模型的输入要求是定长，因此考虑逐帧输入进行预测而不是输入变长的序列去做序列预测。由于需要结合上下文语音帧的信息来确定目标语音帧的表情参数，为了达到引入上下文语音帧的目的，输入不再是一帧，而是对输入的待处理语音以目标语音帧为中心加窗，将窗中的语音帧一起作为短序列输入，此时，窗中的语音帧包括目标语音帧和多个语音帧，构成语音帧集合。可见，通过选取固定的窗长，就能够自然地满足 DNN 模型的输入要求。

30 在这种情况下，语音帧集合中语音帧的数量是由窗长决定的，而窗长反映神经网络映射模型的输入要求，即语音帧集合中语音帧的数量是根据神经网络映射模型确定的。

以图 6 为例，若 PPG 通过 218 维向量表示，窗长为 7，则输入 DNN 模型的输入为包括目标语音帧的 7 个语音帧（有限窗长的 PPG 输入），即语音帧集合中语音帧的数量为 7 帧，每一个语音帧对应一个 PPG（共 218 维向量），7 帧则对应 218×7 维向量，图 6 中每个圆圈表示一维参数。当输入 7 个语音帧分别对应的 PPG 时，便可输出目标语音帧对应的表情参数。

DNN 模型具有建模较为简单, 训练时间较短, 以及可以支持流式工作, 也就是可以逐帧输入而不需要每次输入一整个序列的优势。

需要说明的是, 当神经网络映射模型是 DNN 模型的情况下, 本实施例可以通过多种方式选取语音帧集合中的多个语音帧。其中, 一种可能的实现方式是将目标语音帧的相邻
5 上下文语音帧作为目标语音帧的上下文语音帧。例如, 以目标语音帧为中心, 选取相同数量的相邻上文语音帧和相邻下文语音帧, 参见图 7 所示, 若窗长为 7, 目标语音帧为 X_t , X_t 表示第 t 个语音帧, 则 X_t 的相邻上下文语音帧包括 X_{t-3} 、 X_{t-2} 、 X_{t-1} 、 X_{t+1} 、 X_{t+2} 和 X_{t+3} 。

另一种可能的实现方式是将目标语音帧的间隔上下文语音帧作为目标语音帧的上下文
10 语音帧。本实施例对上下文语音帧的间隔方式不做限定, 例如, 可以采用倍增选帧法, 即按照等比数列的形式倍增地选取上下文语音帧; 也可以按照等差数列的形式倍增地选取上下文语音帧等等。参见图 8 所示, 若窗长为 7, 目标语音帧为 X_t , 按照等比数列的形式倍增地选取上下文语音帧, 则 X_t 的间隔上下文语音帧包括 X_{t-4} 、 X_{t-2} 、 X_{t-1} 、 X_{t+1} 、 X_{t+2} 和 X_{t+4} 。

若神经网络映射模型是 LSTM 模型或 BLSTM 模型, 二者的输入类似, 都可以直接输入
15 表示一句话的语音帧, 而确定表示一句话的语音帧的方式可以是对待处理语音进行语音切分, 例如根据待处理语音中的静音片段进行语音切分, 得到语音切分结果。语音切分结果中切分得到的每个语音片段可以表示一句话, 该语音片段中所包括语音帧的 PPG 可以作为 LSTM 模型或 BLSTM 模型的输入。在这种情况下, 语音切分结果中包括目标语音帧的
20 语音片段中的语音帧可以作为语音帧集合, 此时, 语音帧集合中语音帧的数量是语音切分得到的包括目标语音帧的语音片段中语音帧的数量, 即语音帧集合中语音帧的数量是根据待处理语音的语音切分结果确定的。

其中, LSTM 模型可以参见图 9 所示, 图 9 中每个圆圈表示一维参数。当输入语音帧
25 集合中语音帧分别对应的 PPG 时 (即有限窗长的 PPG 输入), 便可输出目标语音帧对应的表情参数。LSTM 模型的优势是在于能够方便对序列建模, 并且能够捕捉到上下文语音帧的信息, 但更侧重于上文语音帧的信息。

BLSTM 模型可以参见图 10 所示, BLSTM 模型与 LSTM 模型类似, 区别在于 BLSTM
30 中每个隐含层单元能接受序列上下文语音帧两个方向的输入信息, 因此, 相比 LSTM 模型的优势是对下文语音帧的信息也能较好地捕捉到, 更加适合于上下文语音帧都对表情参数的确定有显著影响的情况。

可以理解的是, 在根据语音帧集合中语音帧分别对应的语言学信息确定目标语音帧对
应的表情参数时, 每个目标语音帧对应一个表情参数, 多个目标语音帧对应的表情参数之
间可能存在突变或衔接不连续的情况。为此, 可以对确定出的表情参数进行平滑处理, 避
35 免由于表情参数发生突变, 进而使得根据表情参数驱动动画形象做出的表情连续性更强, 提高动画形象做出表情的真实性。

- 10 -

本申请实施例提供两种平滑处理方法，第一种平滑处理方法是均值平滑。

由于根据语音帧集合中语音帧分别对应的语言学信息，可以确定出语音帧集合中语音帧分别对应的待定表情参数（即语音帧集合中每个语音帧的表情参数）。由于目标语音帧可以出现在不同的语音帧集合中，从而得到多个目标语音帧的待定表情参数，故，基于目标语音帧在不同语音帧集合中分别确定的待定表情参数可以对目标语音帧的表情参数进行平滑处理，计算目标语音帧对应的表情参数。

例如，当目标语音帧为 X_t ，语音帧集合为 $\{X_{t-2}, X_{t-1}, X_t, X_{t+1}, X_{t+2}\}$ ，确定出该语音帧集合中语音帧分别对应的待定表情参数依次是 $\{Y_{t-2}, Y_{t-1}, Y_t, Y_{t+1}, Y_{t+2}\}$ ，目标语音帧 X_t 还可以出现在其它语音帧集合中，例如语音帧集合 $\{X_{t-4}, X_{t-3}, X_{t-2}, X_{t-1}, X_t\}$ 、语音帧集合 $\{X_{t-3}, X_{t-2}, X_{t-1}, X_t, X_{t+1}\}$ 、语音帧集合 $\{X_{t-1}, X_t, X_{t+1}, X_{t+2}, X_{t+3}\}$ 、语音帧集合 $\{X_t, X_{t+1}, X_{t+2}, X_{t+3}, X_{t+4}\}$ ，根据这些语音帧集合，可以确定出目标语音帧在这些集合中分别对应的待定表情参数 Y_t ，即一共得到 5 个目标语音帧 X_t 的待定表情参数。将这 5 个待定表情参数取平均，便可以计算得到目标语音帧 X_t 对应的表情参数。

第二种平滑处理方法是极大似然参数生成算法（Maximum likelihood parameter generation, MLPG）。

由于根据语音帧集合中语音帧分别对应的语言学信息可以确定出语音帧集合中目标语音帧对应的待定表情参数（即目标语音帧的表情参数），同时还可以确定该待定表情参数的一阶差分（或者还有二阶差分），在给定静态参数（待定表情参数）和一阶差分（或者还有二阶差分）的情况下还原出一个似然最大的序列，通过引入的差分的方式来修正待定表情参数的变化从而达到平滑的效果。

在得到平滑后的表情参数之后，就可以通过自然状态下的动画形象，使得动画形象做出待处理语音对应的表情，通过修改动画形象的参数设置，动画形象做出的表情与待处理语音同步。

接下来，将结合具体应用场景，对本申请实施例提供的语音驱动动画方法进行介绍。

在该应用场景中，动画形象用于新闻播报，比如待处理语音为“观众朋友们，大家晚上好，欢迎收看今天的新闻联播”，那么，该动画形象需要做出与该待处理语音对应的口型，使得观众感觉该待处理语音就是该动画形象发出的，增强真实感。为此，在获取待处理语音“观众朋友们，大家晚上好，欢迎收看今天的新闻联播”后，可以确定该待处理语音中语音帧对应的语言学信息。针对待处理语音中的每个语音帧例如目标语音帧，根据语言学信息确定语音帧集合中语音帧分别对应的待定表情参数，将目标语音帧在不同语音帧集合中分别确定的待定表情参数取平均，计算目标语音帧对应的表情参数。这样，可以得到待处理语音中每个语音帧的表情参数，从而驱动动画形象做出“观众朋友们，大家晚上好，欢迎收看今天的新闻联播”对应的口型。

基于前述实施例提供的方法，本实施例还提供一种基于人工智能的语音驱动动画装置 1100，该装置 1100 部署在音视频处理设备上。参见图 11，所述装置 1100 包括获取单元 1101、第一确定单元 1102、第二确定单元 1103 和驱动单元 1104：

5 所述获取单元 1101，用于获取待处理语音，所述待处理语音包括多个语音帧；

所述第一确定单元 1102，用于确定所述待处理语音中语音帧对应的语言学信息，所述语言学信息用于标识所述待处理语音中语音帧所属音素的分布可能性；

所述第二确定单元 1103，用于根据所述语言学信息确定所述待处理语音中语音帧对应的表情参数；

10 所述驱动单元 1104，用于根据所述表情参数驱动动画形象做出对应所述待处理语音的表情。

在一种可能的实现方式中，目标语音帧为所述待处理语音中的一个语音帧，针对所述目标语音帧，所述第二确定单元 1103，用于：

15 确定所述目标语音帧所处的语音帧集合，所述语音帧集合包括所述目标语音帧和所述目标语音帧的上下文语音帧；

根据所述语音帧集合中语音帧分别对应的语言学信息，确定所述目标语音帧对应的表情参数。

20 在一种可能的实现方式中，所述语音帧集合中语音帧的数量是根据神经网络映射模型确定的，或者，所述语音帧集合中语音帧的数量是根据所述待处理语音的语音切分结果确定的。

在一种可能的实现方式中，所述上下文语音帧为所述目标语音帧的相邻上下文语音帧，或者，所述上下文语音帧为所述目标语音帧的间隔上下文语音帧。

在一种可能的实现方式中，所述第二确定单元 1103，用于：

25 根据所述语音帧集合中语音帧分别对应的语言学信息，确定所述语音帧集合中语音帧分别对应的待定表情参数；

根据所述目标语音帧在不同语音帧集合中分别确定的待定表情参数，计算所述目标语音帧对应的表情参数。

在一种可能的实现方式中，所述语言学信息包括音素后验概率、瓶颈特征和嵌入特征中任意一种或多种的组合。

30 在一种可能的实现方式中，所述第二确定单元 1103，用于：

- 12 -

根据所述语言学信息，通过神经网络映射模型确定所述待处理语音中语音帧对应的表情参数；所述神经网络映射模型包括深度神经网络 DNN 模型、长短期记忆网络 LSTM 模型或双向长短期记忆网络 BLSTM 模型。

在一种可能的实现方式中，所述第二确定单元 1103，用于：

- 5 根据所述语言学信息和所述待处理语音对应的情感向量，确定所述待处理语音中语音帧对应的表情参数。

在一种可能的实现方式中，所述第一确定单元 1102，用于：

确定所述待处理语音中语音帧对应的声学特征；

通过自动语音识别模型确定所述声学特征对应的语言学信息。

- 10 在一种可能的实现方式中，所述自动语音识别模型是根据包括了语音片段和音素对应关系的训练样本训练得到的。

本申请实施例还提供了一种用于语音驱动动画的设备，该设备可以通过语音驱动动画，该设备可以为音视频处理设备。下面结合附图对该设备进行介绍。请参见图 12 所示，本申请实施例提供了一种用于语音驱动动画的设备，该设备可以是终端设备，该终端设备可以
15 为包括手机、平板电脑、个人数字助理（Personal Digital Assistant，简称 PDA）、销售终端（Point of Sales，简称 POS）、车载电脑等任意智能终端，以终端设备为手机为例：

- 图 12 示出的是与本申请实施例提供的终端设备相关的手机的部分结构的框图。参考图 12，手机包括：射频（Radio Frequency，简称 RF）电路 1210、存储器 1220、输入单元 1230、显示单元 1240、传感器 1250、音频电路 1260、无线保真（wireless fidelity，简称 WiFi）模块 1270、处理器 1280、以及电源 1290 等部件。本领域技术人员可以理解，图 12 中示出的手机结构并不构成对手机的限定，可以包括比图示更多或更少的部件，或者组合某些部件，或者不同的部件布置。
20

下面结合图 12 对手机的各个构成部件进行具体的介绍：

- RF 电路 1210 可用于收发信息或通话过程中，信号的接收和发送，特别地，将基站的下行信息接收后，给处理器 1280 处理；另外，将设计上行的数据发送给基站。通常，RF 电路 1210 包括但不限于天线、至少一个放大器、收发信机、耦合器、低噪声放大器（Low Noise Amplifier，简称 LNA）、双工器等。此外，RF 电路 1210 还可以通过无线通信与网络和其他设备通信。上述无线通信可以使用任一通信标准或协议，包括但不限于全球移动通讯系统（Global System of Mobile communication，简称 GSM）、通用分组无线服务（General Packet Radio Service，简称 GPRS）、码分多址（Code Division Multiple Access，简称 CDMA）、
25 宽带码分多址（Wideband Code Division Multiple Access，简称 WCDMA）、长期演进（Long Term Evolution，简称 LTE）、电子邮件、短消息服务（Short Messaging Service，简称 SMS）等。
30

存储器 1220 可用于存储软件程序以及模块，处理器 1280 通过运行存储在存储器 1220 的软件程序以及模块，从而执行手机的各种功能应用以及数据处理。存储器 1220 可主要包括存储程序区和存储数据区，其中，存储程序区可存储操作系统、至少一个功能所需的应用程序（比如声音播放功能、图像播放功能等）等；存储数据区可存储根据手机的使用所创建的数据（比如音频数据、电话本等）等。此外，存储器 1220 可以包括高速随机存取存储器，还可以包括非易失性存储器，例如至少一个磁盘存储器件、闪存器件、或其他易失性固态存储器件。

输入单元 1230 可用于接收输入的数字或字符信息，以及产生与手机的用户设置以及功能控制有关的键信号输入。具体地，输入单元 1230 可包括触控面板 1231 以及其他输入设备 1232。触控面板 1231，也称为触摸屏，可收集用户在其上或附近的触摸操作（比如用户使用手指、触笔等任何适合的物体或附件在触控面板 1231 上或在触控面板 1231 附近的操作），并根据预先设定的程式驱动相应的连接装置。可选的，触控面板 1231 可包括触摸检测装置和触摸控制器两个部分。其中，触摸检测装置检测用户的触摸方位，并检测触摸操作带来的信号，将信号传送给触摸控制器；触摸控制器从触摸检测装置上接收触摸信息，并将它转换成触点坐标，再送给处理器 1280，并能接收处理器 1280 发来的命令并加以执行。此外，可以采用电阻式、电容式、红外线以及表面声波等多种类型实现触控面板 1231。除了触控面板 1231，输入单元 1230 还可以包括其他输入设备 1232。具体地，其他输入设备 1232 可以包括但不限于物理键盘、功能键（比如音量控制按键、开关按键等）、轨迹球、鼠标、操作杆等中的一种或多种。

显示单元 1240 可用于显示由用户输入的信息或提供给用户的信息以及手机的各种菜单。显示单元 1240 可包括显示面板 1241，可选的，可以采用液晶显示器（Liquid Crystal Display，简称 LCD）、有机发光二极管（Organic Light-Emitting Diode，简称 OLED）等形式来配置显示面板 1241。进一步的，触控面板 1231 可覆盖显示面板 1241，当触控面板 1231 检测到在其上或附近的触摸操作后，传送给处理器 1280 以确定触摸事件的类型，随后处理器 1280 根据触摸事件的类型在显示面板 1241 上提供相应的视觉输出。虽然在图 12 中，触控面板 1231 与显示面板 1241 是作为两个独立的部件来实现手机的输入和输入功能，但是在某些实施例中，可以将触控面板 1231 与显示面板 1241 集成而实现手机的输入和输出功能。

手机还可包括至少一种传感器 1250，比如光传感器、运动传感器以及其他传感器。具体地，光传感器可包括环境光传感器及接近传感器，其中，环境光传感器可根据环境光线的明暗来调节显示面板 1241 的亮度，接近传感器可在手机移动到耳边时，关闭显示面板 1241 和/或背光。作为运动传感器的一种，加速计传感器可检测各个方向上（一般为三轴）加速度的大小，静止时可检测出重力的大小及方向，可用于识别手机姿态的应用（比如横竖屏切换、相关游戏、磁力计姿态校准）、振动识别相关功能（比如计步器、敲击）等；至于手机还可配置的陀螺仪、气压计、湿度计、温度计、红外线传感器等其他传感器，在此不再赘述。

音频电路 1260、扬声器 1261，传声器 1262 可提供用户与手机之间的音频接口。音频电路 1260 可将接收到的音频数据转换后的电信号，传输到扬声器 1261，由扬声器 1261 转换为声音信号输出；另一方面，传声器 1262 将收集的声音信号转换为电信号，由音频电路 1260 接收后转换为音频数据，再将音频数据输出处理器 1280 处理后，经 RF 电路 1210 以

5 发送给比如另一手机，或者将音频数据输出至存储器 1220 以便进一步处理。

WiFi 属于短距离无线传输技术，手机通过 WiFi 模块 1270 可以帮助用户收发电子邮件、浏览网页和访问流式媒体等，它为用户提供了无线的宽带互联网访问。虽然图 12 示出了 WiFi 模块 1270，但是可以理解的是，其并不属于手机的必须构成，完全可以根据需要在不改变发明的本质的范围内而省略。

10 处理器 1280 是手机的控制中心，利用各种接口和线路连接整个手机的各个部分，通过运行或执行存储在存储器 1220 内的软件程序和/或模块，以及调用存储在存储器 1220 内的数据，执行手机的各种功能和处理数据，从而对手机进行整体监控。可选的，处理器 1280 可包括一个或多个处理单元；优选的，处理器 1280 可集成应用处理器和调制解调处理器，其中，应用处理器主要处理操作系统、用户界面和应用程序等，调制解调处理器主要处理

15 无线通信。可以理解的是，上述调制解调处理器也可以不集成到处理器 1280 中。

手机还包括给各个部件供电的电源 1290（比如电池），优选的，电源可以通过电源管理系统与处理器 1280 逻辑相连，从而通过电源管理系统实现管理充电、放电、以及功耗管理等功能。

尽管未示出，手机还可以包括摄像头、蓝牙模块等，在此不再赘述。

20 在本实施例中，该终端设备所包括的处理器 1280 还具有以下功能：

获取待处理语音，所述待处理语音包括多个语音帧；

确定所述待处理语音中语音帧对应的语言学信息，所述语言学信息用于标识所述待处理语音中语音帧所属音素的分布可能性；

根据所述语言学信息确定所述待处理语音中语音帧对应的表情参数；

25 根据所述表情参数驱动动画形象做出对应所述待处理语音的表情。

本申请实施例还提供服务器，请参见图 13 所示，图 13 为本申请实施例提供的服务器 1300 的结构图，服务器 1300 可因配置或性能不同而产生比较大的差异，可以包括一个或一个以上中央处理器（Central Processing Units，简称 CPU）1322（例如，一个或一个以上处理器）和存储器 1332，一个或一个以上存储应用程序 1342 或数据 1344 的存储介质 1330

30 （例如一个或一个以上海量存储设备）。其中，存储器 1332 和存储介质 1330 可以是短暂存储或持久存储。存储在存储介质 1330 的程序可以包括一个或一个以上模块（图示没标出），每个模块可以包括对服务器中的一系列指令操作。更进一步地，中央处理器 1322 可以设置为与存储介质 1330 通信，在服务器 1300 上执行存储介质 1330 中的一系列指令操作。

服务器 1300 还可以包括一个或一个以上电源 1326，一个或一个以上有线或无线网络接口 1350，一个或一个以上输入输出接口 1358，和/或，一个或一个以上操作系统 1341，例如 Windows Server™，Mac OS X™，Unix™，Linux™，FreeBSD™ 等等。

上述实施例中由服务器所执行的步骤可以基于该图 13 所示的服务器结构。

5 本申请实施例还提供一种计算机可读存储介质，所述计算机可读存储介质用于存储程序代码，所述程序代码用于执行前述各个实施例所述的语音驱动动画方法。

本申请实施例还提供一种包括指令的计算机程序产品，当其在计算机上运行时，使得计算机执行前述各个实施例所述的语音驱动动画方法。

10 本申请的说明书及上述附图中的术语“第一”、“第二”、“第三”、“第四”等（如果存在）是用于区别类似的对象，而不必用于描述特定的顺序或先后次序。应该理解这样使用的数据在适当情况下可以互换，以便这里描述的本申请的实施例例如能够以除了在这里图示或描述的那些以外的顺序实施。此外，术语“包括”和“具有”以及他们的任何变形，意图在于覆盖不排除他的包含，例如，包含了一系列步骤或单元的过程、方法、系统、产品或设备不必限于清楚地列出的那些步骤或单元，而是可包括没有清楚地列出的或对于这些过程、方法、产品或设备固有的其它步骤或单元。

15 应当理解，在本申请中，“至少一个（项）”是指一个或者多个，“多个”是指两个或两个以上。“和/或”，用于描述关联对象的关联关系，表示可以存在三种关系，例如，“A 和/或 B”可以表示：只存在 A，只存在 B 以及同时存在 A 和 B 三种情况，其中 A，B 可以是单数或者复数。字符“/”一般表示前后关联对象是一种“或”的关系。“以下至少一项(个)”或其类似表达，是指这些项中的任意组合，包括单项（个）或复数项（个）的任意组合。例如，a，b 或 c 中的至少一项（个），可以表示：a，b，c，“a 和 b”，“a 和 c”，“b 和 c”，或“a 和 b 和 c”，其中 a，b，c 可以是单个，也可以是多个。

20 在本申请所提供的几个实施例中，应该理解到，所揭露的系统，装置和方法，可以通过其它的方式实现。例如，以上所描述的装置实施例仅仅是示意性的，例如，所述单元的划分，仅仅为一种逻辑功能划分，实际实现时可以有另外的划分方式，例如多个单元或组件可以结合或者可以集成到另一个系统，或一些特征可以忽略，或不执行。另一点，所显示或讨论的相互之间的耦合或直接耦合或通信连接可以是通过一些接口，装置或单元的间接耦合或通信连接，可以是电性，机械或其它的形式。

25 所述作为分离部件说明的单元可以是或者也可以不是物理上分开的，作为单元显示的部件可以是或者也可以不是物理单元，即可以位于一个地方，或者也可以分布到多个网络单元上。可以根据实际的需要选择其中的部分或者全部单元来实现本实施例方案的目的。

30 另外，在本申请各个实施例中的各功能单元可以集成在一个处理单元中，也可以是各个单元单独物理存在，也可以两个或两个以上单元集成在一个单元中。上述集成的单元既可以采用硬件的形式实现，也可以采用软件功能单元的形式实现。

— 16 —

所述集成的单元如果以软件功能单元的形式实现并作为独立的产品销售或使用，可以存储在一个计算机可读取存储介质中。基于这样的理解，本申请的技术方案本质上或者说对现有技术做出贡献的部分或者该技术方案的全部或部分可以以软件产品的形式体现出来，该计算机软件产品存储在一个存储介质中，包括若干指令用以使得一台计算机设备（可以是个人计算机，服务器，或者网络设备等）执行本申请各个实施例所述方法的全部或部分步骤。而前述的存储介质包括：U 盘、移动硬盘、只读存储器（Read-Only Memory，简称 ROM）、随机存取存储器（Random Access Memory，简称 RAM）、磁碟或者光盘等各种可以存储程序代码的介质。

以上所述，以上实施例仅用以说明本申请的技术方案，而非对其限制；尽管参照前述实施例对本申请进行了详细的说明，本领域的普通技术人员应当理解：其依然可以对前述各实施例所记载的技术方案进行修改，或者对其中部分技术特征进行等同替换；而这些修改或者替换，并不使相应技术方案的本质脱离本申请各实施例技术方案的精神和范围。

权 利 要 求

1、一种语音驱动动画方法，所述方法由音视频处理设备执行，所述方法包括：

获取待处理语音，所述待处理语音包括多个语音帧；

5 确定所述待处理语音中语音帧对应的语言学信息，所述语言学信息用于标识所述待处理语音中语音帧所属音素的分布可能性；

根据所述语言学信息确定所述待处理语音中语音帧对应的表情参数；

根据所述表情参数驱动动画形象做出对应所述待处理语音的表情。

10 2、根据权利要求 1 所述的方法，目标语音帧为所述待处理语音中的一个语音帧，针对所述目标语音帧，所述根据所述语言学信息确定所述待处理语音中语音帧对应的表情参数，包括：

确定所述目标语音帧所处的语音帧集合，所述语音帧集合包括所述目标语音帧和所述目标语音帧的上下文语音帧；

根据所述语音帧集合中语音帧分别对应的语言学信息，确定所述目标语音帧对应的表情参数。

15 3、根据权利要求 2 所述的方法，所述语音帧集合中语音帧的数量是根据神经网络映射模型确定的，或者，所述语音帧集合中语音帧的数量是根据所述待处理语音的语音切分结果确定的。

4、根据权利要求 2 所述的方法，所述上下文语音帧为所述目标语音帧的相邻上下文语音帧，或者，所述上下文语音帧为所述目标语音帧的间隔上下文语音帧。

20 5、根据权利要求 2 所述的方法，所述根据所述语音帧集合中语音帧分别对应的语言学信息，确定所述目标语音帧对应的表情参数，包括：

根据所述语音帧集合中语音帧分别对应的语言学信息，确定所述语音帧集合中语音帧分别对应的待定表情参数；

25 根据所述目标语音帧在不同语音帧集合中分别确定的待定表情参数，计算所述目标语音帧对应的表情参数。

6、根据权利要求 1-5 任意一项所述的方法，所述语言学信息包括音素后验概率、瓶颈特征和嵌入特征中任意一种或多种的组合。

7、根据权利要求 1-5 任意一项所述的方法，所述根据所述语言学信息确定所述待处理语音中语音帧对应的表情参数，包括：

— 18 —

根据所述语言学信息，通过神经网络映射模型确定所述待处理语音中语音帧对应的表情参数；所述神经网络映射模型包括深度神经网络 DNN 模型、长短期记忆网络 LSTM 模型或双向长短期记忆网络 BLSTM 模型。

8、根据权利要求 1-5 任意一项所述的方法，所述根据所述语言学信息确定所述待处理语音中语音帧对应的表情参数，包括：

根据所述语言学信息和所述待处理语音对应的情感向量，确定所述待处理语音中语音帧对应的表情参数。

9、根据权利要求 1-5 任意一项所述的方法，所述确定所述待处理语音中语音帧对应的语言学信息，包括：

10 确定所述待处理语音中语音帧对应的声学特征；

通过自动语音识别模型确定所述声学特征对应的语言学信息。

10、根据权利要求 9 所述的方法，所述自动语音识别模型是根据包括了语音片段和音素对应关系的训练样本训练得到的。

11、一种语音驱动动画装置，所述装置部署在音视频处理设备上，所述装置包括获取单元、第一确定单元、第二确定单元和驱动单元：

所述获取单元，用于获取待处理语音，所述待处理语音包括多个语音帧；

所述第一确定单元，用于确定所述待处理语音中语音帧对应的语言学信息，所述语言学信息用于标识所述待处理语音中语音帧所属音素的分布可能性；

20 所述第二确定单元，用于根据所述语言学信息确定所述待处理语音中语音帧对应的表情参数；

所述驱动单元，用于根据所述表情参数驱动动画形象做出对应所述待处理语音的表情。

12、根据权利要求 11 所述的装置，目标语音帧为所述待处理语音中的一个语音帧，针对所述目标语音帧，所述第二确定单元，用于：

25 确定所述目标语音帧所处的语音帧集合，所述语音帧集合包括所述目标语音帧和所述目标语音帧的上下文语音帧；

根据所述语音帧集合中语音帧分别对应的语言学信息，确定所述目标语音帧对应的表情参数。

30 13、根据权利要求 11 所述的装置，所述语音帧集合中语音帧的数量是根据神经网络映射模型确定的，或者，所述语音帧集合中语音帧的数量是根据所述待处理语音的语音切分结果确定的。

14、一种用于语音驱动动画的设备，所述设备包括处理器以及存储器：

— 19 —

所述存储器用于存储程序代码，并将所述程序代码传输给所述处理器；

所述处理器用于根据所述程序代码中的指令执行权利要求 1-10 任意一项所述的方法。

15、一种计算机可读存储介质，所述计算机可读存储介质用于存储程序代码，所述程序代码用于执行权利要求 1-10 任意一项所述的方法。

5 16、一种计算机程序产品，当所述计算机程序产品被执行时，用于执行权利要求 1-10 任意一项所述的方法。

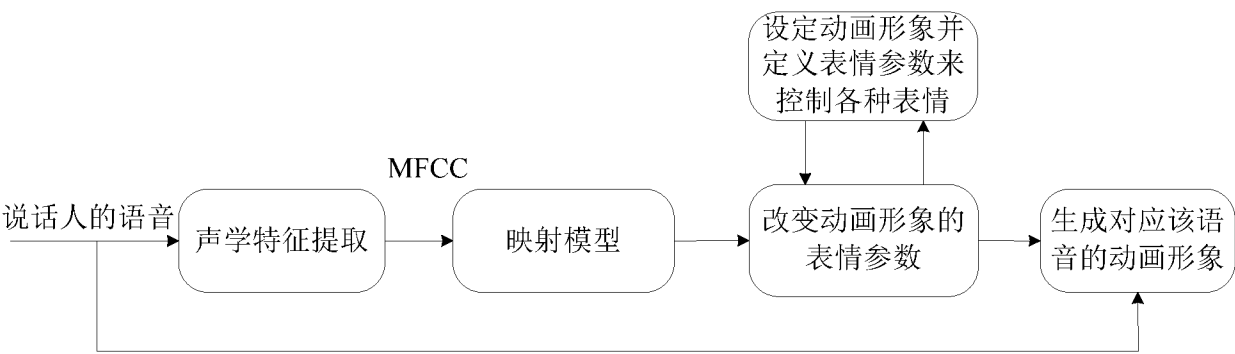


图 1

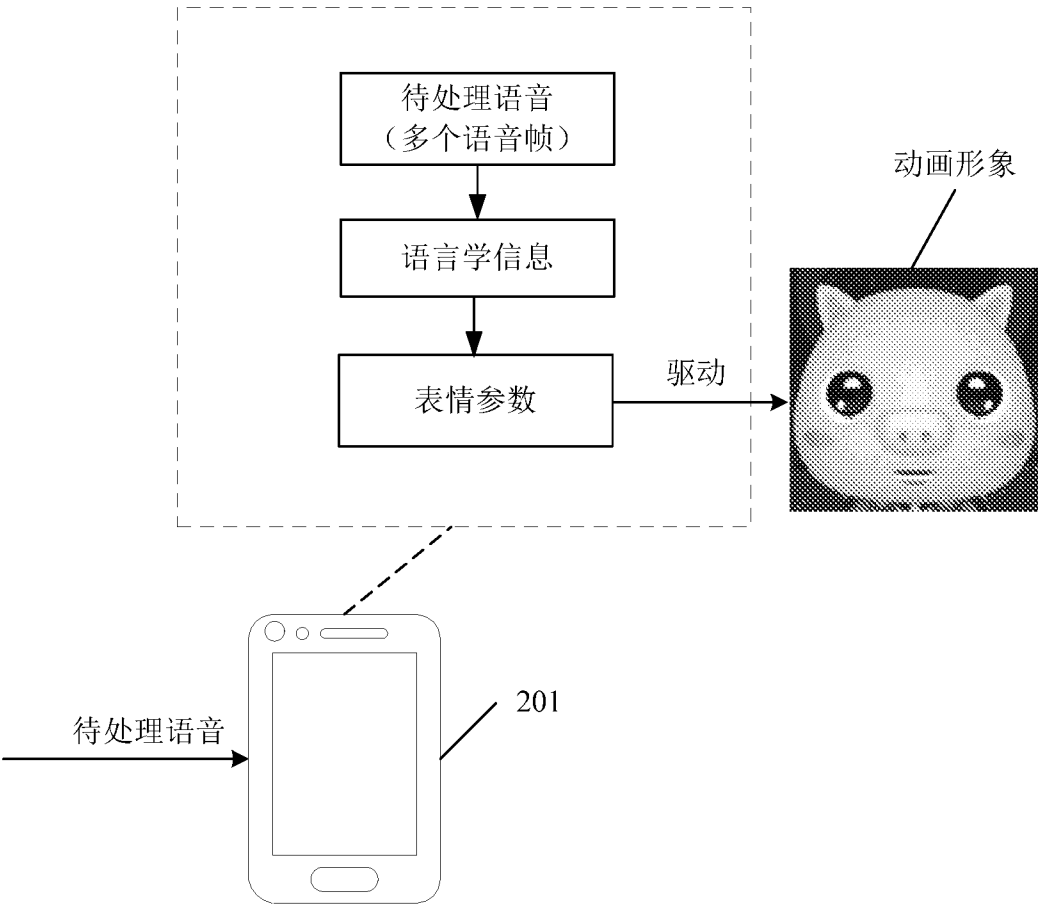


图 2

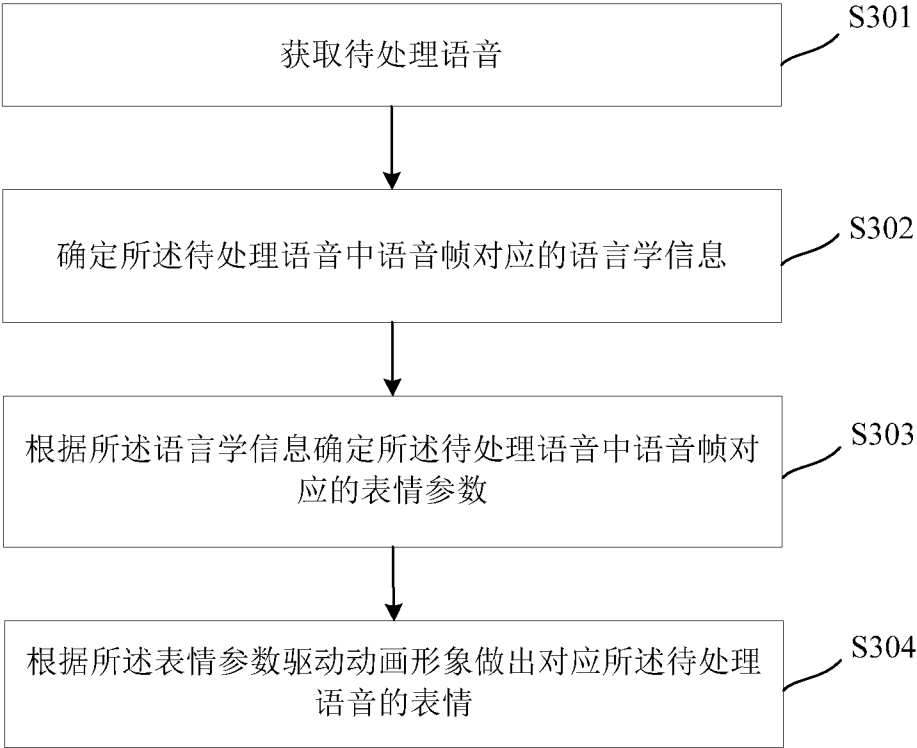


图 3

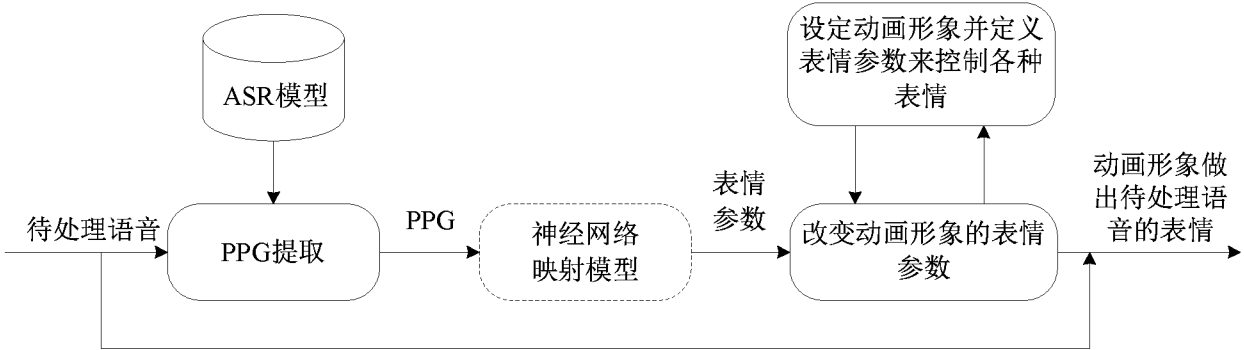


图 4

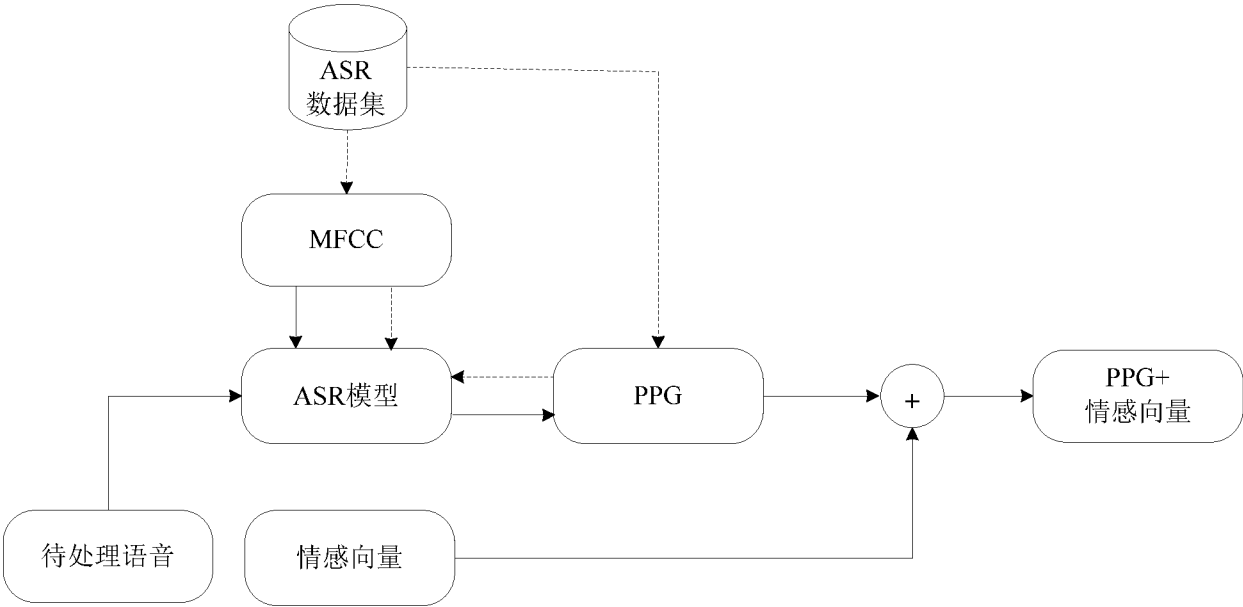


图 5

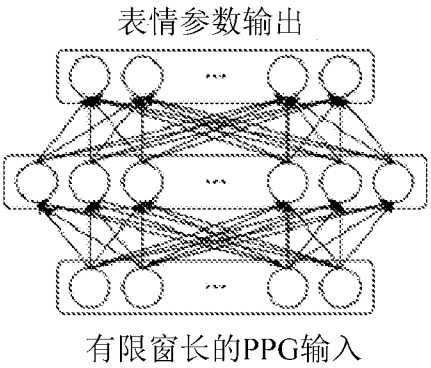


图 6

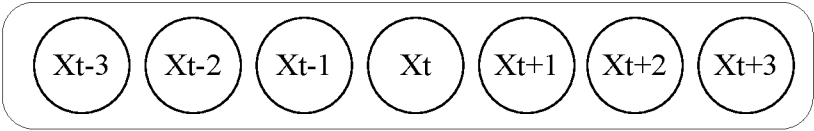


图 7

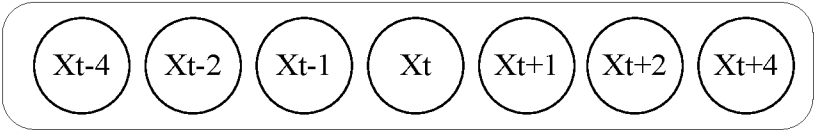


图 8

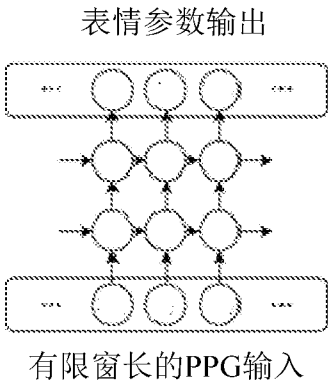


图 9

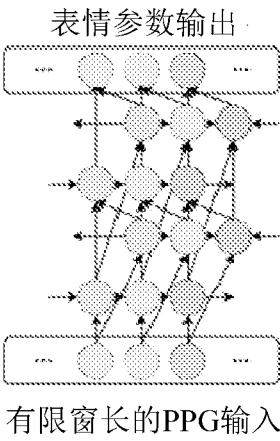


图 10

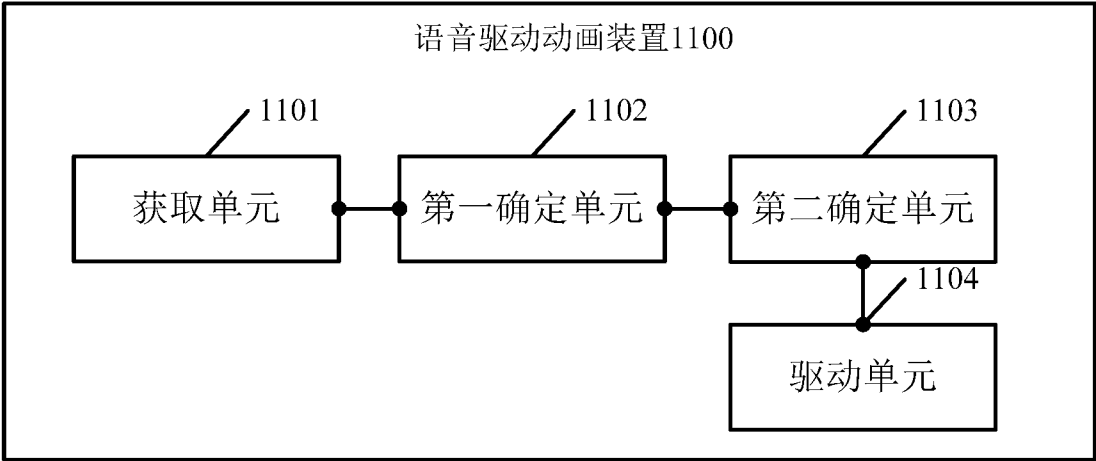


图 11

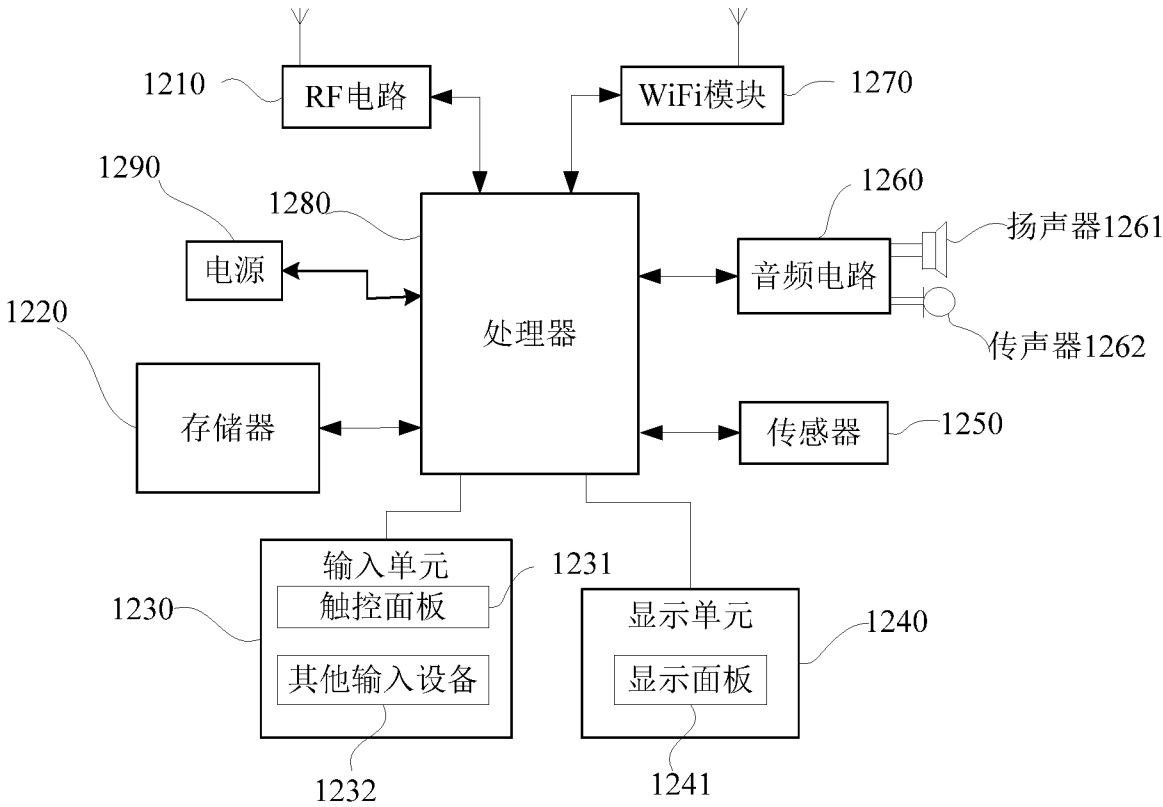


图 12

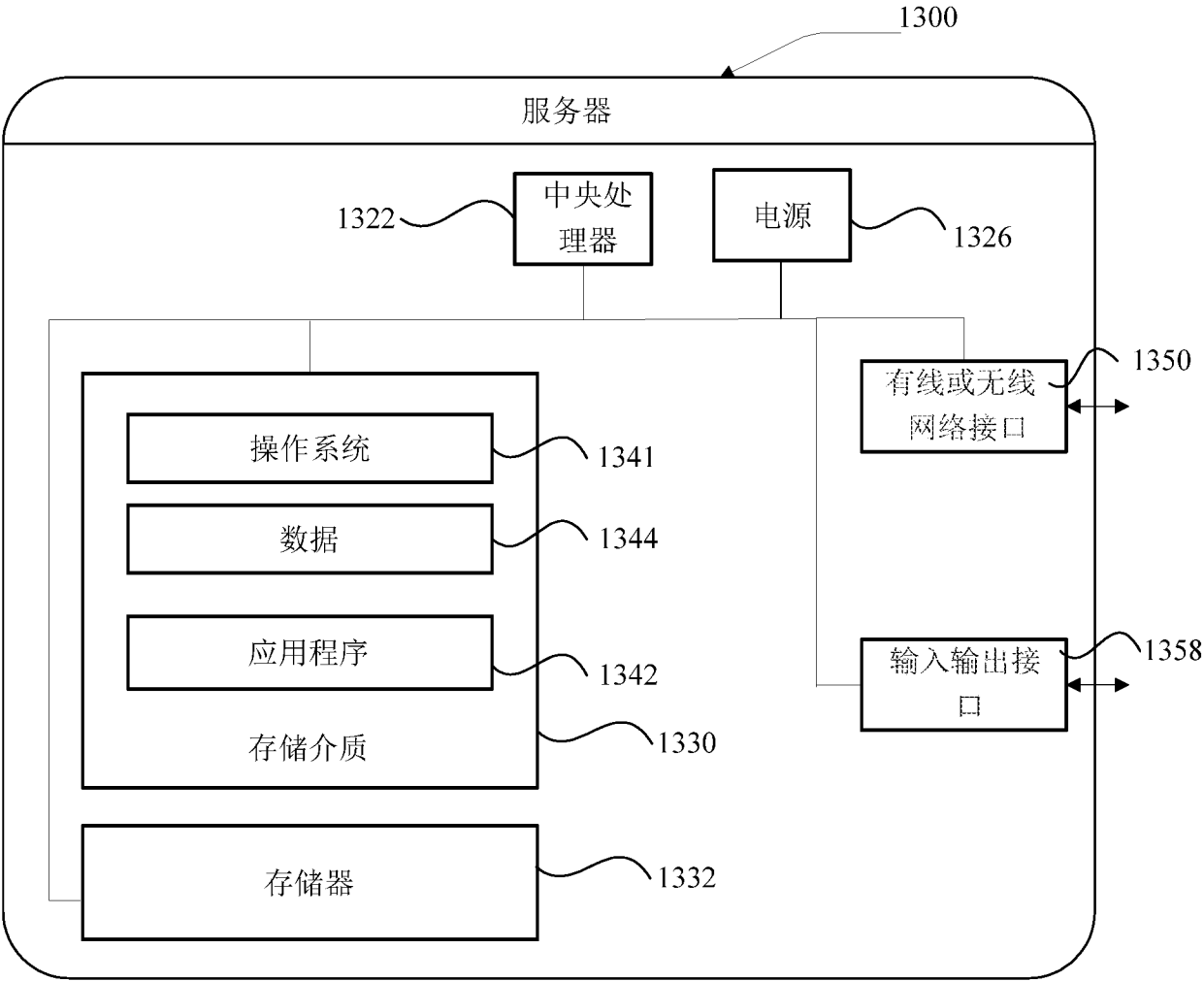


图 13

INTERNATIONAL SEARCH REPORT

International application No.

PCT/CN2020/105046

A. CLASSIFICATION OF SUBJECT MATTER

G10L 15/02(2006.01)i; G10L 15/22(2006.01)i; G10L 15/25(2013.01)i

According to International Patent Classification (IPC) or to both national classification and IPC

B. FIELDS SEARCHED

Minimum documentation searched (classification system followed by classification symbols)

G10L, G06T

Documentation searched other than minimum documentation to the extent that such documents are included in the fields searched

Electronic data base consulted during the international search (name of data base and, where practicable, search terms used)

CNABS, CNTXT, VEN, USTXTOTXT, CNKI, IEEE, 音素, 概率, 表情, 口型, 动画, 帧, 深度, 神经网络, voice, speech, phoneme, face, animation, expression, mouth, shape, DNN, LSTM, BLSTM

C. DOCUMENTS CONSIDERED TO BE RELEVANT

Category*	Citation of document, with indication, where appropriate, of the relevant passages	Relevant to claim No.
PX	CN 110503942 A (TENCENT TECHNOLOGY SHENZHEN CO., LTD.) 26 November 2019 (2019-11-26) claims 1-15, description, paragraphs [0001]-[0149], figures 1-13	1-16
X	CN 108447474 A (BEIJING LINGBAN FUTURE TECHNOLOGY CO., LTD.) 24 August 2018 (2018-08-24) description, paragraphs 29-53	1-16
X	CN 109377540 A (NETEASE (HANGZHOU) NETWORK CO., LTD.) 22 February 2019 (2019-02-22) description, paragraphs 13-21	1-16
A	CN 110009716 A (NETEASE (HANGZHOU) NETWORK CO., LTD.) 12 July 2019 (2019-07-12) entire document	1-16
A	CN 106653052 A (TCL CORP.) 10 May 2017 (2017-05-10) entire document	1-16
A	US 2002024519 A1 (ADAMSOFT CORP.) 28 February 2002 (2002-02-28) entire document	1-16



Further documents are listed in the continuation of Box C.



See patent family annex.

* Special categories of cited documents:

“A” document defining the general state of the art which is not considered to be of particular relevance

“E” earlier application or patent but published on or after the international filing date

“L” document which may throw doubts on priority claim(s) or which is cited to establish the publication date of another citation or other special reason (as specified)

“O” document referring to an oral disclosure, use, exhibition or other means

“P” document published prior to the international filing date but later than the priority date claimed

“T” later document published after the international filing date or priority date and not in conflict with the application but cited to understand the principle or theory underlying the invention

“X” document of particular relevance; the claimed invention cannot be considered novel or cannot be considered to involve an inventive step when the document is taken alone

“Y” document of particular relevance; the claimed invention cannot be considered to involve an inventive step when the document is combined with one or more other such documents, such combination being obvious to a person skilled in the art

“&” document member of the same patent family

Date of the actual completion of the international search

28 October 2020

Date of mailing of the international search report

06 November 2020

Name and mailing address of the ISA/CN

China National Intellectual Property Administration (ISA/
CN)
No. 6, Xitucheng Road, Jimenqiao Haidian District, Beijing
100088
China

Facsimile No. (86-10)62019451

Authorized officer

Telephone No.

INTERNATIONAL SEARCH REPORT

International application No.

PCT/CN2020/105046

C. DOCUMENTS CONSIDERED TO BE RELEVANT

Category*	Citation of document, with indication, where appropriate, of the relevant passages	Relevant to claim No.
A	JP 2006065683 A (KYOCERA COMMUNICATION SYSTEM KK.) 09 March 2006 (2006-03-09) entire document	1-16

INTERNATIONAL SEARCH REPORT
Information on patent family members

International application No.

PCT/CN2020/105046

Patent document cited in search report			Publication date (day/month/year)	Patent family member(s)			Publication date (day/month/year)
CN	110503942	A	26 November 2019	None			
CN	108447474	A	24 August 2018	None			
CN	109377540	A	22 February 2019	None			
CN	110009716	A	12 July 2019	None			
CN	106653052	A	10 May 2017	None			
US	2002024519	A1	28 February 2002	KR	20020022504	A	27 March 2002

国际检索报告

国际申请号

PCT/CN2020/105046

A. 主题的分类

G10L 15/02(2006.01)i; G10L 15/22(2006.01)i; G10L 15/25(2013.01)i

按照国际专利分类(IPC)或者同时按照国家分类和IPC两种分类

B. 检索领域

检索的最低限度文献(标明分类系统和分类号)

G10L, G06T

包含在检索领域中的除最低限度文献以外的检索文献

在国际检索时查阅的电子数据库(数据库的名称, 和使用的检索词(如使用))

CNABS, CNTXT, VEN, USTXTOTXT, CNKI, IEEE, 音素, 概率, 表情, 口型, 动画, 帧, 深度, 神经网络, voice, speech, phoneme, face, animation, expression, mouth, shape, DNN, LSTM, BLSTM

C. 相关文件

类 型*	引用文件, 必要时, 指明相关段落	相关的权利要求
PX	CN 110503942 A (腾讯科技深圳有限公司) 2019年 11月 26日 (2019 - 11 - 26) 权利要求1-15, 说明书第[0001]-[0149]段, 附图1-13	1-16
X	CN 108447474 A (北京灵伴未来科技有限公司) 2018年 8月 24日 (2018 - 08 - 24) 说明书第29-53段	1-16
X	CN 109377540 A (网易杭州网络有限公司) 2019年 2月 22日 (2019 - 02 - 22) 说明书第13-21段	1-16
A	CN 110009716 A (网易杭州网络有限公司) 2019年 7月 12日 (2019 - 07 - 12) 全文	1-16
A	CN 106653052 A (TCL集团股份有限公司) 2017年 5月 10日 (2017 - 05 - 10) 全文	1-16
A	US 2002024519 A1 (ADAMSOFT CORP.) 2002年 2月 28日 (2002 - 02 - 28) 全文	1-16
A	JP 2006065683 A (KYOCERA COMMUNICATION SYSTEM KK.) 2006年 3月 9日 (2006 - 03 - 09) 全文	1-16

☐ 其余文件在C栏的续页中列出。☒ 见同族专利附件。

* 引用文件的具体类型:

“A” 认为不特别相关的表示了现有技术一般状态的文件

“E” 在国际申请日的当天或之后公布的在先申请或专利

“L” 可能对优先权要求构成怀疑的文件, 或为确定另一篇引用文件的公布日而引用的或者因其他特殊理由而引用的文件(如具体说明的)

“O” 涉及口头公开、使用、展览或其他方式公开的文件

“P” 公布日先于国际申请日但迟于所要求的优先权日的文件

“T” 在申请日或优先权日之后公布, 与申请不相抵触, 但为了理解发明之理论或原理的在后文件

“X” 特别相关的文件, 单独考虑该文件, 认定要求保护的发明不是新颖的或不具有创造性

“Y” 特别相关的文件, 当该文件与另一篇或者多篇该类文件结合并且这种结合对于本领域技术人员为显而易见时, 要求保护的发明不具有创造性

“&” 同族专利的文件

国际检索实际完成的日期

2020年 10月 28日

国际检索报告邮寄日期

2020年 11月 6日

ISA/CN的名称和邮寄地址

中国国家知识产权局(ISA/CN)

中国北京市海淀区蓟门桥西土城路6号 100088

传真号 (86-10)62019451

授权官员

吴琼

电话号码 (86-10)62089510

国际检索报告
关于同族专利的信息

国际申请号
PCT/CN2020/105046

检索报告引用的专利文件			公布日 (年/月/日)	同族专利			公布日 (年/月/日)
CN	110503942	A	2019年 11月 26日	无			
CN	108447474	A	2018年 8月 24日	无			
CN	109377540	A	2019年 2月 22日	无			
CN	110009716	A	2019年 7月 12日	无			
CN	106653052	A	2017年 5月 10日	无			
US	2002024519	A1	2002年 2月 28日	KR	20020022504	A	2002年 3月 27日