

(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



(10) 国际公布号

WO 2020/237869 A1

(43) 国际公布日

2020 年 12 月 3 日 (03.12.2020)

(51) 国际专利分类号:

G06F 16/33 (2019.01) G06F 16/332 (2019.01)

(21) 国际申请号:

PCT/CN2019/102922

(22) 国际申请日:

2019 年 8 月 28 日 (28.08.2019)

(25) 申请语言:

中文

(26) 公布语言:

中文

(30) 优先权:

201910467185.X 2019年5月31日 (31.05.2019) CN

(71) 申请人: 平安科技(深圳)有限公司(PING AN TECHNOLOGY (SHENZHEN) CO., LTD.) [CN/CN];

中国广东省深圳市福田区福田街道福安社区益田路 5033 号平安金融中心 23 楼, Guangdong 518000 (CN)。

(72) 发明人: 王健宗(WANG, Jianzong); 中国广东省深圳市福田区福田街道福安社区益田路 5033 号平安金融中心 23 楼, Guangdong 518000 (CN)。 孙奥兰(SUN, Aolan); 中国广东省深圳市福田区福田街道福安社区益田路 5033 号平安金融中心 23 楼, Guangdong 518000 (CN)。 彭话易(PENG, Huayi); 中国广东省深圳市福田区福田街道福安社区益田路 5033 号平安金融中心 23 楼, Guangdong 518000 (CN)。 程宁(CHENG, Ning); 中国广东省深圳市福田区福田街道福安社区益田路 5033 号平安金融中心 23 楼, Guangdong 518000 (CN)。

(74) 代理人: 深圳众鼎专利商标代理事务所(普通合伙)(SHENZHEN ZHONGDING INTELLECTUAL PROPERTY AGENCY); 中国广东省深圳市罗湖区笋岗东路 3012 号中民时代广场 B 座 7 楼 701 室, Guangdong 518000 (CN)。

(54) Title: QUESTION INTENTION RECOGNITION METHOD AND APPARATUS, COMPUTER DEVICE, AND STORAGE MEDIUM

(54) 发明名称: 一种问题意图识别方法、装置、计算机设备及存储介质

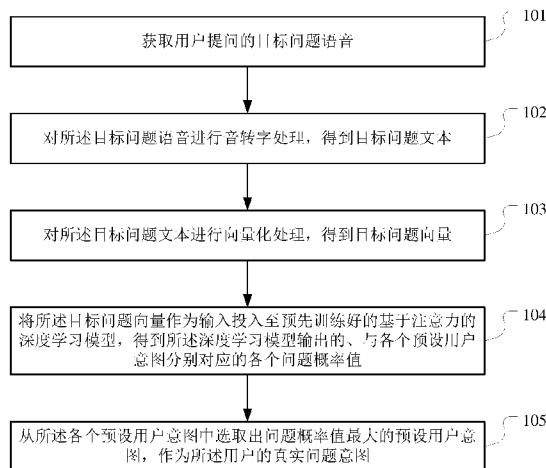


图 2

- 101 Obtain a target question voice of a user asking a question
- 102 Perform voice-to-character conversion processing on the target question voice to obtain a target question text
- 103 Perform vectorization processing on the target question text to obtain a target question vector
- 104 Use the target question vector as an input and input same into a pre-trained attention-based deep learning model so as to obtain each question probability value that is outputted by the deep learning model and separately corresponds to each preset user intention
- 105 Select the preset user intention having a maximum question probability value from various preset user intentions to be used as the real question intention of the user

(57) Abstract: Disclosed in the present application are a question intention recognition method and apparatus, a computer device, and a storage medium, being applied to the technical field of deep learning and used for solving a problem that the real intention of a question of a user is hard to understand by means of an existing technical means. The method provided by the present application comprises: obtaining a target question voice of a user asking a question; performing voice-to-character conversion processing on the target question voice to obtain a target question text; performing vectorization processing on the target question text to obtain a target question vector; using the target question vector as an input and inputting same into a pre-trained attention-based deep learning model so as to obtain each question probability value that is outputted by the deep learning model and separately corresponds to each preset user intention, wherein each question probability value separately represents a probability that the target question text belongs to the preset user intention corresponding to each question probability value; and selecting the preset user intention having a maximum question probability value from various preset user intentions to be used as the real question intention of the user.

(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, 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, SM, ST, SV, SY, TH, TJ, TM, TN, TR, TT, TZ, UA, UG, US, UZ, VC, VN, 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))。

(57) 摘要: 本申请公开了一种问题意图识别方法、装置、计算机设备及存储介质, 应用于深度学习技术领域, 用于解决现有技术手段难以理解用户问题的真实意图的问题。本申请提供的方法包括: 获取用户提问的目标问题语音; 对目标问题语音进行音转字处理, 得到目标问题文本; 对目标问题文本进行向量化处理, 得到目标问题向量; 将目标问题向量作为输入投入至预先训练好的基于注意力的深度学习模型, 得到深度学习模型输出的、与各个预设用户意图分别对应的各个问题概率值, 每个问题概率值各自表征了目标问题文本属于与每个问题概率值对应的预设用户意图的概率; 从各个预设用户意图中选取出问题概率值最大的预设用户意图, 作为用户的真实问题意图。

一种问题意图识别方法、装置、计算机设备及存储介质

本申请以 2019 年 05 月 31 日提交的申请号为 201910467185.X，名称为“一种问题意图识别方法、装置、计算机设备及存储介质”的中国发明专利申请为基础，并要求其优先权。

技术领域

本申请涉及深度学习技术领域，尤其涉及一种问题意图识别方法、装置、计算机设备及存储介质。

背景技术

近年来，语音助手的使用需求逐渐增加，用户可以通过向语音助手提问得到反馈的答案，帮助用户解决疑问。然而，现有语音助手仅能从用户问题的字面意思出发来搜索答案，常常出现答非所问的情况，难以满足用户的需求。

因此，寻找一种能够准确理解用户问题意图的方法成为本领域技术人员亟需解决的问题。

发明内容

本申请实施例提供一种问题意图识别方法、装置、计算机设备及存储介质，以解决现有技术手段难以理解用户问题的真实意图的问题。

一种问题意图识别方法，其特征在于，包括：

获取用户提问的目标问题语音；

对所述目标问题语音进行音转字处理，得到目标问题文本；

对所述目标问题文本进行向量化处理，得到目标问题向量；

将所述目标问题向量作为输入投入至预先训练好的基于注意力的深度学习模型，得到所述深度学习模型输出的、与各个预设用户意图分别对应的各个问题概率值，每个问题概率值各自表征了所述目标问题文本属于与所述每个问题概率值对应的预设用户意图的概率；

从所述各个预设用户意图中选取出问题概率值最大的预设用户意图，作为所述用户的真实问题意图。

一种问题意图识别装置，其特征在于，包括：

问题语音获取模块，用于获取用户提问的目标问题语音；

音转字模块，用于对所述目标问题语音进行音转字处理，得到目标问题文本；

文本向量化模块，用于对所述目标问题文本进行向量化处理，得到目标问题向量；

问题识别模块，用于将所述目标问题向量作为输入投入至预先训练好的基于注意力的深度学习模型，得到所述深度学习模型输出的、与各个预设用户意图分别对应的各个问题概率值，每个问题概率值各自表征了所述目标问题文本属于与所述每个问题概率值对应的预设用户意图的概率；

真实意图选取模块，用于从所述各个预设用户意图中选取出问题概率值最大的预设用户意图，作为所述用户的真实问题意图。

一种计算机设备，包括存储器、处理器以及存储在所述存储器中并可在所述处理器上运行的计算机可读指令，所述处理器执行所述计算机可读指令时实现上述问题意图识别方法的步骤。

一个或多个存储有计算机可读指令的可读存储介质，所述计算机可读存储介质存储

有计算机可读指令，使得所述一个或多个处理器执行上述问题意图识别方法的步骤。

本申请的一个或多个实施例的细节在下面的附图和描述中提出，本申请的其他特征和优点将从说明书、附图以及权利要求变得明显。

附图说明

为了更清楚地说明本申请实施例的技术方案，下面将对本申请实施例的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 是本申请一实施例中问题意图识别方法的一应用环境示意图；

图 2 是本申请一实施例中问题意图识别方法的一流程图；

图 3 是本申请一实施例中问题意图识别方法步骤 103 在一个应用场景下的流程示意图；

图 4 是本申请一实施例中问题意图识别方法在一个应用场景下预先训练深度学习模型的流程示意图；

图 5 是本申请一实施例中问题意图识别方法步骤 104 在一个应用场景下的流程示意图；

图 6 是本申请一实施例中问题意图识别方法在一个应用场景下提供答案给用户的流程示意图；

图 7 是本申请一实施例中问题意图识别装置在一个应用场景下的结构示意图；

图 8 是本申请一实施例中问题意图识别装置在另一个应用场景下的结构示意图；

图 9 是本申请一实施例中问题识别模块的结构示意图；

图 10 是本申请一实施例中计算机设备的一示意图。

具体实施方式

下面将结合本申请实施例中的附图，对本申请实施例中的技术方案进行清楚、完整地描述，显然，所描述的实施例是本申请一部分实施例，而不是全部的实施例。基于本申请中的实施例，本领域普通技术人员在没有作出创造性劳动前提下所获得的所有其他实施例，都属于本申请保护的范围。

本申请提供的问题意图识别方法，可应用在如图 1 的应用环境中，其中，客户端通过网络与服务器进行通信。其中，该客户端可以但不限于各种个人计算机、笔记本电脑、智能手机、平板电脑和便携式可穿戴设备。服务器可以用独立的服务器或者是多个服务器组成的服务器集群来实现。

在一实施例中，如图 2 所示，提供一种问题意图识别方法，以该方法应用在图 1 中的服务器为例进行说明，包括如下步骤：

101、获取用户提问的目标问题语音；

本实施例中，服务器可以根据实际使用的需要或者应用场景的需要获取用户提问的目标问题语音。例如，服务器可以与客户端通信连接，该客户端提供给某场所内的用户咨询问题，用户通过客户端的麦克风输入语音，客户端将该语音上传给服务器，从而服务器获取到的该语音即为目标问题语音。或者，服务器也可以执行对大批量的话术录音识别用户意图的任务，某数据库预先收集大量的来自用户提问的话术录音，然后通过网络将这些话术录音传输给服务器，从而服务器获取到的这些话术录音即为用户提问的目标问题语音。

可以理解的是，服务器还可以通过多种方式获取到用户提问的目标问题语音，对此不再过多赘述。

需要说明的是，本实施例中所说的目标问题语音一般是指用户提问时采集的声音数

据。

102、对所述目标问题语音进行音转字处理，得到目标问题文本；

容易理解的是，服务器在获取到目标问题语音之后，可以将该目标问题语音转换为文字，即可得到目标问题文本。比如，服务器可以采用 ASR（Automatic Speech Recognition）自动语音识别技术对该目标问题语音进行识别，完成音转字处理，得到该目标问题文本。

103、对所述目标问题文本进行向量化处理，得到目标问题向量；

在得到目标问题文本之后，为了便于后续深度学习模型的识别，服务器需要对该目标问题文本进行向量化处理，即将目标问题文本转化为向量的方式表示，从而得到目标问题向量。具体地，服务器可以将目标问题文本以数据矩阵的形式记载，在数据矩阵中，目标问题文本中的每个字词映射为该数据矩阵中的一个行向量。

为便于理解，如图 3 所示，进一步地，步骤 103 可以包括：

201、将所述目标问题文本中的各个目标字词分别转换为 GloVe（Global Vectors for Word Representation）词向量，得到初始问题向量；

202、判断所述各个目标字词是否均被 GloVe 词向量覆盖，若是，则执行步骤 203，若否，则执行步骤 204；

203、确定所述初始问题向量为目标问题向量；

204、将未被 GloVe 词向量覆盖的目标字词转换为 TransE 词向量，得到补充向量；

205、将所述补充向量添加至所述初始问题向量，得到目标问题向量。

对于上述步骤 201，GloVe 的全称叫 Global Vectors for Word Representation，它是现有一个基于全局词频统计（count-based & overall statistics）的词表征（word representation）工具，它可以把一个单词表达成一个由实数组成的向量。本实施例中，服务器使用 GloVe 将所述目标问题文本中的各个目标字词分别转换成词向量，从而得到初始问题向量。

对于上述步骤 202，考虑到用户提出的问题中可能包含有专有名词，例如姓名、地点等，这些专有名词难以被 GloVe 全覆盖。因此，服务器可以判断所述各个目标字词是否均被 GloVe 词向量覆盖，如果该目标问题文本中的各个目标字词均已被覆盖，则可以执行步骤 203，直接确定所述初始问题向量为目标问题向量；反之，如果该目标问题文本中的各个目标字词未被全覆盖，则需要执行后续步骤 204 和步骤 205。

对于上述步骤 203，由上述内容可知，若所述各个目标字词均被 GloVe 词向量覆盖，则服务器可以确定所述初始问题向量为目标问题向量。

对于上述步骤 204，若所述各个目标字词中任一个目标字词未被 GloVe 词向量覆盖，则可知该目标问题文本中存在无法被 GloVe 词向量覆盖的目标字词，为了补充这一部分的缺失，服务器可以将未被 GloVe 词向量覆盖的目标字词转换为 TransE 词向量，得到补充向量。

需要说明的，TransE，又称知识库方法，是一种现有的有效学习专有名词的算法模型，可以将学习到的字词转换为分布式向量表示。本实施例中，服务器可以采用 TransE 将未被 GloVe 词向量覆盖的目标字词进行向量转换，得到补充向量。

对于步骤 205，可知，服务器在得到补充向量之后，可以使用该补充向量添加至该初始问题向量中，以填补初始问题向量的缺失，从而得到该目标问题文本对应的目标问题向量。举例说明，假设该目标问题文本为“小明吃饭吗”，该句子中包括“小明”、“吃饭”和“吗”三个目标字词。服务器使用 GloVe 将“吃饭”和“吗”转换为词向量，分别为[1234]和[1235]，对于“小明”一词，服务器使用 TransE 将其转换为词向量，得到[1236]，再将[1236]添加至[1234]和[1235]中，得到该目标问题向量为[1236]、[1234]和[1235]，其中，该目标问题向量可以以一维向量的形式表达，即[1236]、[1234]和[1235] 合并为[123612341235]作为该目标问题向量，也可以以二维向量的形式表达，即[1236]、[1234]和[1235]分别作为一个二维向

量的行向量，得到目标向量为：
$$\begin{bmatrix} 1 & 2 & 3 & 6 \\ 1 & 2 & 3 & 4 \\ 1 & 2 & 3 & 5 \end{bmatrix}$$
。

104、将所述目标问题向量作为输入投入至预先训练好的基于注意力的深度学习模型，得到所述深度学习模型输出的、与各个预设用户意图分别对应的各个问题概率值，每个问题概率值各自表征了所述目标问题文本属于与所述每个问题概率值对应的预设用户意图的概率；

可以理解的是，服务器在得到目标问题向量之后，可以将所述目标问题向量作为输入投入至预先训练好的基于注意力的深度学习模型，得到所述深度学习模型输出的、与各个预设用户意图分别对应的各个问题概率值。其中，每个问题概率值各自表征了所述目标问题文本属于与所述每个问题概率值对应的预设用户意图的概率。可以理解的是，这些问题概率值与各个预设用户意图一一对应，当某个问题概率值越大，则代表了用户提问的问题属于该问题概率值对应的预设用户意图的可能性越高。

为便于理解，下面将对基于注意力的深度学习模型的训练过程进行详细描述。如图 4 所示，进一步地，所述深度学习模型包括第一双层循环神经网络、Key-Value 记忆网络、第二双层循环神经网络和内容相似度计算网络四个部分，所述深度学习模型通过以下步骤预先训练好：

301、收集属于所述各个预设用户意图的样本问题语音；

302、对收集到的样本问题语音分别进行音转字处理，得到样本问题文本；

303、对所述样本问题文本进行向量化处理，得到样本问题向量；

304、针对每个样本问题向量，为所述每个样本问题向量分别针对各个预设用户意图设定标记值，得到所述每个样本问题向量的各个标记值，其中，与所述每个样本问题向量对应的预设用户意图的标记值最大；

305、将所有所述样本问题向量分别输入第一双层循环神经网络进行编码，得到各个样本问题向量各自对应的语境特征向量；

306、将各个所述语境特征向量分别投入 Key-Value 记忆网络中进行迭代计算，直到达到预设的迭代次数或模型收敛，然后输出迭代计算后的各个语境特征向量，所述 Key-Value 记忆网络中参与迭代计算各个 Key 值为各个样本问题文本中主语和谓语对应的向量；

307、将迭代计算后的各个语境特征向量分别输入第二双层循环神经网络进行解码，得到各个样本结果向量；

308、针对每个样本结果向量，通过内容相似度计算网络计算所述每个样本结果向量与各个意图向量之间的相似度，得到所述每个样本结果向量对应的各个相似度值，作为样本概率值，所述各个意图向量是指所述各个预设用户意图向量化后的向量值；

309、以所述每个样本结果向量对应的各个样本概率值为调整目标，调整所述第一双层循环神经网络、所述 Key-Value 记忆网络、所述第二双层循环神经网络和所述内容相似度计算网络的网络参数，以最小化所述每个样本结果向量对应的各个样本概率值与各个目标标记值之间的误差，所述各个目标标记值是指所述每个样本结果向量对应的样本问题向量的各个标记值；

310、若每个样本结果向量对应的各个样本概率值与各个目标标记值之间的误差满足预设的训练终止条件，则确定所述深度学习模型已训练好。

对于步骤 301，本实施例中，针对实际应用场景，工作人员可以预先在服务器上设置好需要训练的各个预设用户意图，例如可以包括“问好”、“希望挂断”、“委婉拒绝”等意图，针对这些预设用户意图，工作人员还需要在具体应用场景下收集各自对应的用户话术，即用户的语音作为样本问题语音，比如用户实际咨询的问题语音。在收集样本问题语音时，

服务器可以通过专业咨询平台、网络客服等渠道收集属于各个预设用户意图的样本问题语音。需要说明的是，每个预设用户意图对应的样本问题语音应当达到一定的数量级，各个预设用户意图之间样本问题语音的数量可以有一定差距，但不应相差过远，避免影响对深度学习模型的训练效果。例如，可以收集到的样本问题语音为：“问好”对应的样本问题语音的数量为 100 万条，“希望挂断”对应的样本问题语音的数量为 20 万条，“委婉拒绝”对应的样本问题语音的数量为 30 万条。

对于步骤 302，与上述步骤 102 同理，服务器可以对收集到的样本问题语音分别进行音转字处理，得到样本问题文本，此处不再赘述。

对于步骤 303，与上述步骤 103 同理，服务器也可以对所述样本问题文本进行向量化处理，得到样本问题向量。

对于步骤 304，可以理解的是，在训练之前，需要对样本问题向量进行标记，本实施例中由于需要针对多个预设用户意图进行训练，因此应当针对不同的预设用户意图分别进行设定标记值。举例说明，假设共 3 个预设用户意图，分别为“问好”、“希望挂断”和“委婉拒绝”，假设共 100 万个样本问题向量，针对 1 号样本问题向量，由于该样本问题向量对应的样本问题文本的真实意图为“问好”，则将 1 号样本问题向量对“问好”的标记值设为 1，对“希望挂断”和“委婉拒绝”的标记值均设为 0；针对 2 号样本问题向量，由于该样本问题向量对应的样本问题文本的真实意图为“希望挂断”，则将 2 号样本问题向量对“希望挂断”的标记值设为 1，对“问好”和“委婉拒绝”的标记值均设为 0；同理将所有 100 万个样本问题向量均分别针对各个预设用户意图设定标记值，即完成训练前的样本标注工作。

需要说明的是，上述举例中将样本问题向量对应的预设用户意图的标记值记为 1，其它的预设用户意图的标记值记为 0，这只是其中一种标记值的设定方式。比如，也可以将样本问题向量对应的预设用户意图的标记值记为 0.9，其它的预设用户意图的标记值记为 0.8、0.7、0.6 等等，只要比 0.9 小即可，保证样本问题向量对应的预设用户意图的标记值在所有标记值中最大即可。

对于步骤 305，本实施例中，该深度学习模型中设有第一双层循环神经网络作为编码器，第二双层循环神经网络作为解码器，其中，该第一双层循环神经网络即为两层的 RNN。步骤 305 中，首先，将样本问题向量输入到第一层 RNN 中进行卷积计算，通过合理设置第一层 RNN 的卷积核，计算出初步的特征表示（向量），然后将初步的特征表示再输入第二层 RNN 做卷积计算，完成对该初步的特征表示的特征遍历，计算得到的特征表示结果可以认为是该样本问题向量对应的语境特征向量。步骤 305 的卷积计算过程由于完成了对样本问题向量的特征提取和编码，因此可以认为该第一双层循环神经网络为深度学习模型中的编码器。

对于步骤 306，本实施例中，预先将各个样本问题向量的主语、谓语和宾语存储为 Key-Value 记忆网络中的各个 Key-Value 对。其中，Key 值由样本问题向量中主语和谓语对应的部分向量组成，表示为（主语，谓语），Value 值由样本问题向量中宾语对应的部分向量组成，示为（宾语）。需要说明的是，上述所说的主语、谓语和宾语均以向量的形式表示，具体是指一个样本问题文本中主语、谓语和宾语对应字词的词向量。例如，对于“小明吃饭吗”这一样本问题文本，主语为“小明”一词对应的向量[1236]，谓语为“吃饭”一词对应的向量[1234]，从而该样本问题文本的 Key 值可以表示为([1236], [1234])。

具体地，步骤 306 中，服务器是将每个语境特征向量与 Key-Value 记忆网络中各个 Key-Value 对的 Key 值进行相似度计算，并将计算得到的相似度值存储为所述每个语境特征向量与各个 Key 值之间的注意力权重，然后利用该注意力权重来更新所述每个语境特征向量，如此反复迭代，直到达到预设的迭代次数或模型收敛，服务器再输出经过多次迭代计算后的语境特征向量。

为便于理解，上述过程可以通过下述公式一进行表达：

$$q_{j+1} = R_j(q_j + \sum_i^N \text{Soft max}(q_j A \Phi_K))$$

其中, q_j 表示第 j 次迭代计算时语境特征向量, q_{j+1} 表示第 $j+1$ 次迭代计算时语境特征向量, i 各个 Key-Value 对序号, 也等于所述各个样本问题向量的数量, Φ_K 为表示 Key 值, A 和 R_j 为 Key-Value 记忆力网络的网络参数, Softmax 函数为归一化指数函数, 具体为

$$\sigma(z)_j = \frac{e^{z_j}}{\sum_{k=1}^K e^{z_k}}。$$

需要说明的是, 本实施例中预设的迭代次数具体可以根据实际情况设定, 比如可以设定为 10 次。而关于 Key-Value 记忆网络何时达到模型收敛的判断, 本实施例中可以采用现有的损失函数 (Loss function) 来进行判定, 比如可以采用互熵损失函数 (Softmax Loss) 来实现对 Key-Value 记忆网络的模型收敛判断, 此处不再展开描述。

对于步骤 307, 可以理解的是, 服务器执行不住 306 之后, 得到迭代计算后的各个语境特征向量, 此时, 还需要将这些语境特征向量输入到第二双层循环神经网络中进行解码。本实施例中, 该第二双层循环神经网络与上述的第一双层循环神经网络原理类似, 同样是设置两层 RNN, 每层 RNN 通过合理设置卷积核对所述各个语境特征向量进行卷积计算, 完成对各个语境特征向量的解码操作, 从而最后得到各个样本结果向量。可以理解的是, 服务器得到的样本结果向量的维度、尺寸与上述样本问题向量的维度、尺寸一致。

对于步骤 308, 服务器上预设了需要训练的各个预设用户意图, 比如“问好”、“希望挂断”、“委婉拒绝”等意图, 在训练时, 为了衡量这些预设用户意图与样本结果向量之间的相似程度, 服务器需要将各个预设用户意图向量化, 得到各个意图向量。服务器在得到各个样本结果向量之后, 可以针对每个样本结果向量, 通过内容相似度计算网络计算所述每个样本结果向量与各个意图向量之间的相似度, 得到所述每个样本结果向量对应的各个相似度值, 并将这里的相似度值看作输出结果与预设用户意图之间的对应的可能性, 记为样本概率值。

具体地, 步骤 308 可以通过如下公式计算所述每个样本结果向量与各个意图向量之间的相似度, 即该内容相似度计算网络可以通过以下公式进行表达:

$$\alpha_{i,m} = \exp(e_{i,m}) / \sum_{m=1}^M \exp(e_{i,m})$$

$$\exp(e_{i,m}) = \omega^T \tanh(Ws_{m-1} + Vh_m + b)$$

其中, ω 和 b 均为向量, W 和 V 为矩阵, ω 、 b 、 W 、 V 为内容相似度计算网络的网路参数, $\alpha_{i,m}$ 表示第二双层循环神经网络输出的样本结果向量 s_{m-1} 与意图向量 h_m 之间的相似度值, 也即样本概率值, m 为意图向量的序号, M 为所述各个意图向量的数量, 即共 M 个预设用户意图。

对于步骤 309, 可以理解的是, 在训练深度学习模型的过程中, 需要调整该深度学习模型的参数, 具体在本实施例中, 即调整所述第一双层循环神经网络、所述 Key-Value 记忆网络、所述第二双层循环神经网络和所述内容相似度计算网络的网路参数, 比如上述的

ω 、 b 、 W 、 V 、 A 和 R_j ，等等。通过调整这些网络参数可以影响该深度学习模型最终的输出结果，使得每个样本结果向量对应的各个样本概率值与各个目标标记值之间的误差最小化。

对于上述步骤 310，在调节上述网络参数的过程中，可以判断每个样本结果向量对应的各个样本概率值与各个目标标记值之间的误差满足预设的训练终止条件，若满足，则说明该深度学习模型中的各个网络参数已经调整到位，可以确定该深度学习模型已训练完成；反之，若不满足，则说明该深度学习模型还需要继续训练。

其中，该训练终止条件可以根据实际使用情况预先设定，具体地，可以将该训练终止条件设定为：若每个样本结果向量对应的各个样本概率值与各个目标标记值之间的误差均小于指定误差值，则认为其满足该预设的训练终止条件。或者，也可以将其设为：使用验证集中的样本问题语音执行上述步骤 301-309，若深度学习模型输出的样本概率值与标记值之间的误差在一定范围内，则认为其满足该预设的训练终止条件。其中，该验证集中的样本问题语音的收集与上述步骤 301 类似，具体地，可以执行上述步骤 301 收集得到各个预设用户意图的样本问题语音后，将收集得到的样本问题语音中的一部分划分为训练集，剩余的样本问题语音划分为验证集。比如，可以将收集得到的样本问题语音中随机划分 80% 作为后续训练深度学习模型的训练集的样本，将其它的 20% 划分为后续验证深度学习模型是否训练完成，也即是否满足预设训练终止条件的验证集的样本。

上面描述了基于注意力的深度学习模型的预先训练过程，为便于理解，下面承接上述训练过程的内容，详细描述一下使用该深度学习模型在实际使用中对目标问题向量的识别过程。如图 5 所示，更进一步地，所述深度学习模型包括第一双层循环神经网络、Key-Value 记忆网络、第二双层循环神经网络和内容相似度计算网络四个部分，步骤 104 可以包括：

401、将所述目标问题向量输入所述第一双层循环神经网络进行编码，得到目标语境特征向量；

402、将所述目标语境特征向量投入所述 Key-Value 记忆网络中进行迭代计算，直到达到预设的迭代次数或模型收敛，然后输出迭代计算后的目标语境特征向量；

403、将迭代计算后的目标语境特征向量输入所述第二双层循环神经网络进行解码，得到目标结果向量；

404、通过所述内容相似度计算网络计算所述目标结果向量与各个意图向量之间的相似度，得到所述目标结果向量对应的各个相似度值，作为各个问题概率值。

上述步骤 401-404 与上述步骤 305-308 的原理类似，此处不再过多赘述。

在步骤 401-404 中，服务器先将所述目标问题向量输入所述第一双层循环神经网络进行编码，得到目标语境特征向量；然后，将所述目标语境特征向量投入所述 Key-Value 记忆网络中进行迭代计算，直到达到预设的迭代次数或模型收敛，然后输出迭代计算后的目标语境特征向量；接着，将迭代计算后的目标语境特征向量输入所述第二双层循环神经网络进行解码，得到目标结果向量；最后，服务器通过所述内容相似度计算网络计算所述目标结果向量与各个意图向量之间的相似度，得到所述目标结果向量对应的各个相似度值，作为各个问题概率值。可知，这里得到的各个问题概率值表征了用户提问的目标问题语音分别属于各个预设用户意图的可能性大小，某个问题概率值越大，则表示该问题概率值对应的预设用户意图越有可能是用户本次提供的真实意图。

105、从所述各个预设用户意图中选取出问题概率值最大的预设用户意图，作为所述用户的真实问题意图。

服务器在得到所述深度学习模型输出的、与各个预设用户意图分别对应的各个问题概率值之后，由于每个问题概率值各自表征了所述目标问题文本属于与所述每个问题概率值对应的预设用户意图的概率，因此，服务器可以从所述各个预设用户意图中选取出问题概

率值最大的预设用户意图，这里选取出的预设用户意图在所有预设用户意图中最有可能是用户提问的真实意图，因此将其确定为所述用户的真实问题意图。

本实施例中，服务器在确定出用户的真实问题意图之后，还可以通过预先设置多个问题答案元组，从这些问题答案元组中选取与真实问题意图对应的答案提供给用户，进而直接解决用户的提问。如图 6 所示，进一步地，在步骤 105 之后，本方法还可以包括：

501、获取预设的各个问题答案元组，每个问题答案元组由一个预设用户意图和与所述一个预设用户意图对应的答案组成，其中，所述各个问题答案元组的预设用户意图各不相同；

502、从所述各个问题答案元组中选取与预设用户意图与所述真实问题意图相同的一个问题答案元组，作为命中元组；

503、将所述命中元组的答案反馈至所述用户。

对于步骤 501，可以理解的是，服务器上可以预先设置各个问题答案元组，每个问题答案元组由一个预设用户意图和与所述一个预设用户意图对应的答案组成，比如，某个问题答案元组为（问好；谢谢，你呢？），其中，“问好”为该元组中的预设用户意图，“谢谢，你呢？”为该元组中的答案。并且，为了让所述各个问题答案元组覆盖所有预设用户意图，因此服务器上的所述各个问题答案元组的预设用户意图各不相同，即针对每个预设用户意图设置一个问题答案元组。

需要说明的是，为了让问题答案元组中的答案更加准确且具有代表性，这些问题答案元组中的答案可以预先经过人工筛选和唯一处理，将针对预设用户意图的最合适答案设置在问题答案元组中。比如，在一个实际应用场景下，可以预设 5.2k 个预设用户意图和 30.8k 个答案。在这 5.2k 预设用户意图中，仅仅保留其中的 330 个对本应用场景有意义的。相似地，30.8k 答案中，人为挑选其中 642 个答案用作设定各个问题答案元组。最后，经过对答案的唯一处理和筛选，可以大约有 1k 个问题答案元组，覆盖了 1k 个用户可能提问的问题的真实意图并给出准确的答案。

对于步骤 502 和步骤 503，容易理解的是，服务器在确定出用户的真实问题意图之后，可以从所述各个问题答案元组中选取与预设用户意图与所述真实问题意图相同的一个问题答案元组，作为命中元组，然后将所述命中元组的答案反馈至所述用户，即可完成对用户提问的解答。

本申请实施例中，首先，获取用户提问的目标问题语音；然后，对所述目标问题语音进行音转字处理，得到目标问题文本；接着，对所述目标问题文本进行向量化处理，得到目标问题向量；再之，将所述目标问题向量作为输入投入至预先训练好的基于注意力的深度学习模型，得到所述深度学习模型输出的、与各个预设用户意图分别对应的各个问题概率值，每个问题概率值各自表征了所述目标问题文本属于与所述每个问题概率值对应的预设用户意图的概率；最后，从所述各个预设用户意图中选取出问题概率值最大的预设用户意图，作为所述用户的真实问题意图。可见，本申请通过基于注意力的深度学习模型可以准确地从用户提问的目标问题语音出发识别出用户的真实意图，提升了意图识别的准确性，当应用于语音助手等语音识别情境时，可以大大减少答非所问的情况出现。

应理解，上述实施例中各步骤的序号的大小并不意味着执行顺序的先后，各过程的执行顺序应以其功能和内在逻辑确定，而不对本申请实施例的实施过程构成任何限定。

在一实施例中，提供一种问题意图识别装置，该问题意图识别装置与上述实施例中问题意图识别方法一一对应。如图 7 所示，该问题意图识别装置包括问题语音获取模块 601、音转字模块 602、文本向量化模块 603、问题识别模块 604 和真实意图选取模块 605。各功能模块详细说明如下：

问题语音获取模块 601, 用于获取用户提问的目标问题语音;

音转字模块 602, 用于对所述目标问题语音进行音转字处理, 得到目标问题文本;

文本向量化模块 603, 用于对所述目标问题文本进行向量化处理, 得到目标问题向量;

问题识别模块 604, 用于将所述目标问题向量作为输入投入至预先训练好的基于注意力的深度学习模型, 得到所述深度学习模型输出的、与各个预设用户意图分别对应的各个问题概率值, 每个问题概率值各自表征了所述目标问题文本属于与所述每个问题概率值对应的预设用户意图的概率;

真实意图选取模块 605, 用于从所述各个预设用户意图中选取出问题概率值最大的预设用户意图, 作为所述用户的真实问题意图。

如图 8 所示, 进一步地, 所述深度学习模型包括第一双层循环神经网络、Key-Value 记忆网络、第二双层循环神经网络和内容相似度计算网络四个部分, 所述深度学习模型可以通过以下模块预先训练好:

样本收集模块 606, 用于收集属于所述各个预设用户意图的样本问题语音;

样本音转字模块 607, 用于对收集到的样本问题语音分别进行音转字处理, 得到样本问题文本;

样本向量化模块 608, 用于对所述样本问题文本进行向量化处理, 得到样本问题向量;

样本标记模块 609, 用于针对每个样本问题向量, 为所述每个样本问题向量分别针对各个预设用户意图设定标记值, 得到所述每个样本问题向量的各个标记值, 其中, 与所述每个样本问题向量对应的预设用户意图的标记值最大;

向量编码模块 610, 用于将所有所述样本问题向量分别输入第一双层循环神经网络进行编码, 得到各个样本问题向量各自对应的语境特征向量;

迭代计算模块 611, 用于将各个所述语境特征向量分别投入 Key-Value 记忆网络中进行迭代计算, 直到达到预设的迭代次数或模型收敛, 然后输出迭代计算后的各个语境特征向量, 所述 Key-Value 记忆网络中参与迭代计算各个 Key 值为各个样本问题文本中主语和谓语对应的向量;

向量解码模块 612, 用于将迭代计算后的各个语境特征向量分别输入第二双层循环神经网络进行解码, 得到各个样本结果向量;

相似度值计算模块 613, 用于针对每个样本结果向量, 通过内容相似度计算网络计算所述每个样本结果向量与各个意图向量之间的相似度, 得到所述每个样本结果向量对应的各个相似度值, 作为样本概率值, 所述各个意图向量是指所述各个预设用户意图向量化后的向量值;

网络参数调整模块 614, 用于以所述每个样本结果向量对应的各个样本概率值为调整目标, 调整所述第一双层循环神经网络、所述 Key-Value 记忆网络、所述第二双层循环神经网络和所述内容相似度计算网络的网络参数, 以最小化所述每个样本结果向量对应的各个样本概率值与各个目标标记值之间的误差, 所述各个目标标记值是指所述每个样本结果向量对应的样本问题向量的各个标记值;

训练完成确定模块 615, 用于若每个样本结果向量对应的各个样本概率值与各个目标标记值之间的误差满足预设的训练终止条件, 则确定所述深度学习模型已训练好。

如图 9 所示, 进一步地, 所述问题识别模块 604 可以包括:

编码单元 6041, 用于将所述目标问题向量输入所述第一双层循环神经网络进行编码, 得到目标语境特征向量;

向量迭代计算单元 6042, 用于将所述目标语境特征向量投入所述 Key-Value 记忆网络中进行迭代计算, 直到达到预设的迭代次数或模型收敛, 然后输出迭代计算后的目标语境特征向量;

解码单元 6043, 用于将迭代计算后的目标语境特征向量输入所述第二双层循环神经

网络进行解码，得到目标结果向量；

相似度计算单元 6044，用于通过所述内容相似度计算网络计算所述目标结果向量与各个意图向量之间的相似度，得到所述目标结果向量对应的各个相似度值，作为各个问题概率值。

进一步地，所述文本向量化模块可以包括：

第一转换单元，用于将所述目标问题文本中的各个目标字词分别转换为 GloVe 词向量，得到初始问题向量；

字词判断单元，用于判断所述各个目标字词是否均被 GloVe 词向量覆盖；

问题向量确定单元，用于若所述字词判断单元的判断结果为是，则确定所述初始问题向量为目标问题向量；

第二转换单元，用于若所述字词判断单元的判断结果为否，则将未被 GloVe 词向量覆盖的目标字词转换为 TransE 词向量，得到补充向量；

向量添加单元，用于将所述补充向量添加至所述初始问题向量，得到目标问题向量。

进一步地，所述问题意图识别装置还可以包括：

答案元组获取模块，用于获取预设的各个问题答案元组，每个问题答案元组由一个预设用户意图和与所述一个预设用户意图对应的答案组成，其中，所述各个问题答案元组的预设用户意图各不相同；

答案元组选取模块，用于从所述各个问题答案元组中选取与预设用户意图与所述真实问题意图相同的一个问题答案元组，作为命中元组；

答案反馈模块，用于将所述命中元组的答案反馈至所述用户。

关于问题意图识别装置的具体限定可以参见上文中对于问题意图识别方法的限定，在此不再赘述。上述问题意图识别装置中的各个模块可全部或部分通过软件、硬件及其组合来实现。上述各模块可以硬件形式内嵌于或独立于计算机设备中的处理器中，也可以以软件形式存储于计算机设备中的存储器中，以便于处理器调用执行以上各个模块对应的操作。

在一个实施例中，提供了一种计算机设备，该计算机设备可以是服务器，其内部结构图可以如图 10 所示。该计算机设备包括通过系统总线连接的处理器、存储器、网络接口和数据库。其中，该计算机设备的处理器用于提供计算和控制能力。该计算机设备的存储器包括可读存储介质、内存储器。该可读存储介质存储有操作系统、计算机可读指令和数据库。该内存储器为可读存储介质中的操作系统和计算机可读指令的运行提供环境。该计算机设备的数据库用于存储问题意图识别方法中涉及到的数据。该计算机设备的网络接口用于与外部的终端通过网络连接通信。该计算机可读指令被处理器执行时以实现一种问题意图识别方法。本实施例所提供的可读存储介质包括非易失性可读存储介质和易失性可读存储介质。

在一个实施例中，提供了一种计算机设备，包括存储器、处理器及存储在存储器上并可在处理器上运行的计算机可读指令，处理器执行计算机可读指令时实现上述实施例中问题意图识别方法的步骤，例如图 2 所示的步骤 101 至步骤 105。或者，处理器执行计算机可读指令时实现上述实施例中问题意图识别装置各模块/单元的功能，例如图 7 所示模块 601 至模块 605 的功能。为避免重复，这里不再赘述。

在一个实施例中，提供了一种计算机可读存储介质，该一个或多个存储有计算机可读指令的可读存储介质，计算机可读指令被一个或多个处理器执行时，使得一个或多个处理器执行计算机可读指令时实现上述方法实施例中问题意图识别方法的步骤，或者，该一个或多个存储有计算机可读指令的可读存储介质，计算机可读指令被一个或多个处理器执行时，使得一个或多个处理器执行计算机可读指令时实现上述装置实施例中问题意图识别装

置中各模块/单元的功能。为避免重复，这里不再赘述。本实施例所提供的可读存储介质包括非易失性可读存储介质和易失性可读存储介质。

本领域普通技术人员可以理解实现上述实施例方法中的全部或部分流程，是可以通过计算机可读指令来指令相关的硬件来完成，所述的计算机可读指令可存储于一计算机可读存储介质中，该计算机可读指令在执行时，可包括如上述各方法的实施例的流程。其中，本申请所提供的各实施例中所使用的对存储器、存储、数据库或其它介质的任何引用，均可包括和/或易失性存储器。存储器可包括只读存储器（ROM）、可编程 ROM（PROM）、电可编程 ROM（EPROM）、电可擦除可编程 ROM（EEPROM）或闪存。易失性存储器可包括随机存取存储器（RAM）或者外部高速缓冲存储器。作为说明而非局限，RAM 以多种形式可得，诸如静态 RAM（SRAM）、动态 RAM（DRAM）、同步 DRAM（SDRAM）、双数据率 SDRAM（DDRSDRAM）、增强型 SDRAM（ESDRAM）、同步链路（Synchlink）DRAM（SLDRAM）、存储器总线（Rambus）直接 RAM（RDRAM）、直接存储器总线动态 RAM（DRDRAM）、以及存储器总线动态 RAM（RDRAM）等。

所属领域的技术人员可以清楚地了解到，为了描述的方便和简洁，仅以上述各功能单元、模块的划分进行举例说明，实际应用中，可以根据需要而将上述功能分配由不同的功能单元、模块完成，即将所述装置的内部结构划分成不同的功能单元或模块，以完成以上描述的全部或者部分功能。

以上所述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，均应包含在本申请的保护范围之内。

权 利 要 求 书

1. 一种问题意图识别方法，其特征在于，包括：

获取用户提问的目标问题语音；

对所述目标问题语音进行音转字处理，得到目标问题文本；

对所述目标问题文本进行向量化处理，得到目标问题向量；

将所述目标问题向量作为输入投入至预先训练好的基于注意力的深度学习模型，得到所述深度学习模型输出的、与各个预设用户意图分别对应的各个问题概率值，每个问题概率值各自表征了所述目标问题文本属于与所述每个问题概率值对应的预设用户意图的概率；

从所述各个预设用户意图中选取出问题概率值最大的预设用户意图，作为所述用户的真实问题意图。

2. 根据权利要求 1 所述的问题意图识别方法，其特征在于，所述深度学习模型包括第一双层循环神经网络、Key-Value 记忆网络、第二双层循环神经网络和内容相似度计算网络四个部分，所述深度学习模型通过以下步骤预先训练好：

收集属于所述各个预设用户意图的样本问题语音；

对收集到的样本问题语音分别进行音转字处理，得到样本问题文本；

对所述样本问题文本进行向量化处理，得到样本问题向量；

针对每个样本问题向量，为所述每个样本问题向量分别针对各个预设用户意图设定标记值，得到所述每个样本问题向量的各个标记值，其中，与所述每个样本问题向量对应的预设用户意图的标记值最大；

将所有所述样本问题向量分别输入第一双层循环神经网络进行编码，得到各个样本问题向量各自对应的语境特征向量；

将各个所述语境特征向量分别投入 Key-Value 记忆网络中进行迭代计算，直到达到预设的迭代次数或模型收敛，然后输出迭代计算后的各个语境特征向量，所述 Key-Value 记忆网络中参与迭代计算各个 Key 值为各个样本问题文本中主语和谓语对应的向量；

将迭代计算后的各个语境特征向量分别输入第二双层循环神经网络进行解码，得到各个样本结果向量；

针对每个样本结果向量，通过内容相似度计算网络计算所述每个样本结果向量与各个意图向量之间的相似度，得到所述每个样本结果向量对应的各个相似度值，作为样本概率值，所述各个意图向量是指所述各个预设用户意图向量化后的向量值；

以所述每个样本结果向量对应的各个样本概率值为调整目标，调整所述第一双层循环神经网络、所述 Key-Value 记忆网络、所述第二双层循环神经网络和所述内容相似度计算网络的网络参数，以最小化所述每个样本结果向量对应的各个样本概率值与各个目标标记值之间的误差，所述各个目标标记值是指所述每个样本结果向量对应的样本问题向量的各个标记值；

若每个样本结果向量对应的各个样本概率值与各个目标标记值之间的误差满足预设的训练终止条件，则确定所述深度学习模型已训练好。

3. 根据权利要求 2 所述的问题意图识别方法，其特征在于，所述将所述目标问题向量作为输入投入至预先训练好的基于注意力的深度学习模型，得到所述深度学习模型输出的、与各个预设用户意图分别对应的各个问题概率值包括：

将所述目标问题向量输入所述第一双层循环神经网络进行编码，得到目标语境特征向量；

将所述目标语境特征向量投入所述 Key-Value 记忆网络中进行迭代计算，直到达到预设的迭代次数或模型收敛，然后输出迭代计算后的目标语境特征向量；

将迭代计算后的目标语境特征向量输入所述第二双层循环神经网络进行解码,得到目标结果向量;

通过所述内容相似度计算网络计算所述目标结果向量与各个意图向量之间的相似度,得到所述目标结果向量对应的各个相似度值,作为各个问题概率值。

4. 根据权利要求 1 所述的问题意图识别方法,其特征在于,所述对所述目标问题文本进行向量化处理,得到目标问题向量包括:

将所述目标问题文本中的各个目标字词分别转换为 GloVe 词向量,得到初始问题向量;

判断所述各个目标字词是否均被 GloVe 词向量覆盖;

若所述各个目标字词均被 GloVe 词向量覆盖,则确定所述初始问题向量为目标问题向量;

若所述各个目标字词中任一目标字词未被 GloVe 词向量覆盖,则将未被 GloVe 词向量覆盖的目标字词转换为 TransE 词向量,得到补充向量;

将所述补充向量添加至所述初始问题向量,得到目标问题向量。

5. 根据权利要求 1 至 4 中任一项所述的问题意图识别方法,其特征在于,在从所述各个预设用户意图中选取出问题概率值最大的预设用户意图,作为所述用户的真实问题意图之后,还包括:

获取预设的各个问题答案元组,每个问题答案元组由一个预设用户意图和与所述一个预设用户意图对应的答案组成,其中,所述各个问题答案元组的预设用户意图各不相同;

从所述各个问题答案元组中选取预设用户意图与所述真实问题意图相同的一个问题答案元组,作为命中元组;

将所述命中元组的答案反馈至所述用户。

6. 一种问题意图识别装置,其特征在于,包括:

问题语音获取模块,用于获取用户提问的目标问题语音;

音转字模块,用于对所述目标问题语音进行音转字处理,得到目标问题文本;

文本向量化模块,用于对所述目标问题文本进行向量化处理,得到目标问题向量;

问题识别模块,用于将所述目标问题向量作为输入投入至预先训练好的基于注意力的深度学习模型,得到所述深度学习模型输出的、与各个预设用户意图分别对应的各个问题概率值,每个问题概率值各自表征了所述目标问题文本属于与所述每个问题概率值对应的预设用户意图的概率;

真实意图选取模块,用于从所述各个预设用户意图中选取出问题概率值最大的预设用户意图,作为所述用户的真实问题意图。

7. 根据权利要求 6 所述的问题意图识别装置,其特征在于,所述深度学习模型包括第一双层循环神经网络、Key-Value 记忆网络、第二双层循环神经网络和内容相似度计算网络四个部分,所述深度学习模型通过以下模块预先训练好:

样本收集模块,用于收集属于所述各个预设用户意图的样本问题语音;

样本音转字模块,用于对收集到的样本问题语音分别进行音转字处理,得到样本问题文本;

样本向量化模块,用于对所述样本问题文本进行向量化处理,得到样本问题向量;

样本标记模块,用于针对每个样本问题向量,为所述每个样本问题向量分别针对各个预设用户意图设定标记值,得到所述每个样本问题向量的各个标记值,其中,与所述每个样本问题向量对应的预设用户意图的标记值最大;

向量编码模块,用于将所有所述样本问题向量分别输入第一双层循环神经网络进行编码,得到各个样本问题向量各自对应的语境特征向量;

迭代计算模块,用于将各个所述语境特征向量分别投入 Key-Value 记忆网络中进行迭

代计算,直到达到预设的迭代次数或模型收敛,然后输出迭代计算后的各个语境特征向量,所述 Key-Value 记忆网络中参与迭代计算各个 Key 值为各个样本问题文本中主语和谓语对应的向量;

向量解码模块,用于将迭代计算后的各个语境特征向量分别输入第二双层循环神经网络进行解码,得到各个样本结果向量;

相似度值计算模块,用于针对每个样本结果向量,通过内容相似度计算网络计算所述每个样本结果向量与各个意图向量之间的相似度,得到所述每个样本结果向量对应的各个相似度值,作为样本概率值,所述各个意图向量是指所述各个预设用户意图向量化后的向量值;

网络参数调整模块,用于以所述每个样本结果向量对应的各个样本概率值为调整目标,调整所述第一双层循环神经网络、所述 Key-Value 记忆网络、所述第二双层循环神经网络和所述内容相似度计算网络的网络参数,以最小化所述每个样本结果向量对应的各个样本概率值与各个目标标记值之间的误差,所述各个目标标记值是指所述每个样本结果向量对应的样本问题向量的各个标记值;

训练完成确定模块,用于若每个样本结果向量对应的各个样本概率值与各个目标标记值之间的误差满足预设的训练终止条件,则确定所述深度学习模型已训练好。

8. 根据权利要求 7 所述的问题意图识别装置,其特征在于,所述问题识别模块包括:

编码单元,用于将所述目标问题向量输入所述第一双层循环神经网络进行编码,得到目标语境特征向量;

向量迭代计算单元,用于将所述目标语境特征向量投入所述 Key-Value 记忆网络中进行迭代计算,直到达到预设的迭代次数或模型收敛,然后输出迭代计算后的目标语境特征向量;

解码单元,用于将迭代计算后的目标语境特征向量输入所述第二双层循环神经网络进行解码,得到目标结果向量;

相似度计算单元,用于通过所述内容相似度计算网络计算所述目标结果向量与各个意图向量之间的相似度,得到所述目标结果向量对应的各个相似度值,作为各个问题概率值。

9. 根据权利要求 6 所述的问题意图识别装置,其特征在于,所述文本向量化模块包括:

第一转换单元,用于将所述目标问题文本中的各个目标字词分别转换为 GloVe 词向量,得到初始问题向量;

字词判断单元,用于判断所述各个目标字词是否均被 GloVe 词向量覆盖;

问题向量确定单元,用于若所述字词判断单元的判断结果为是,则确定所述初始问题向量为目标问题向量;

第二转换单元,用于若所述字词判断单元的判断结果为否,则将未被 GloVe 词向量覆盖的目标字词转换为 TransE 词向量,得到补充向量;

向量添加单元,用于将所述补充向量添加至所述初始问题向量,得到目标问题向量。

10. 根据权利要求 6 至 9 中任一项所述的问题意图识别装置,其特征在于,所述问题意图识别装置还包括:

答案元组获取模块,用于获取预设的各个问题答案元组,每个问题答案元组由一个预设用户意图和与所述一个预设用户意图对应的答案组成,其中,所述各个问题答案元组的预设用户意图各不相同;

答案元组选取模块,用于从所述各个问题答案元组中选取预设用户意图与所述真实问题意图相同的一个问题答案元组,作为命中元组;

答案反馈模块,用于将所述命中元组的答案反馈至所述用户。

11. 一种计算机设备,包括存储器、处理器以及存储在所述存储器中并可在所述处理

器上运行的计算机可读指令，其特征在于，所述处理器执行所述计算机可读指令时实现如下步骤：

获取用户提问的目标问题语音；

对所述目标问题语音进行音转字处理，得到目标问题文本；

对所述目标问题文本进行向量化处理，得到目标问题向量；

将所述目标问题向量作为输入投入至预先训练好的基于注意力的深度学习模型，得到所述深度学习模型输出的、与各个预设用户意图分别对应的各个问题概率值，每个问题概率值各自表征了所述目标问题文本属于与所述每个问题概率值对应的预设用户意图的概率；

从所述各个预设用户意图中选取出问题概率值最大的预设用户意图，作为所述用户的真实问题意图。

12. 根据权利要求 11 所述的计算机设备，其特征在于，所述深度学习模型包括第一双层循环神经网络、Key-Value 记忆网络、第二双层循环神经网络和内容相似度计算网络四个部分，所述深度学习模型通过以下步骤预先训练好：

收集属于所述各个预设用户意图的样本问题语音；

对收集到的样本问题语音分别进行音转字处理，得到样本问题文本；

对所述样本问题文本进行向量化处理，得到样本问题向量；

针对每个样本问题向量，为所述每个样本问题向量分别针对各个预设用户意图设定标记值，得到所述每个样本问题向量的各个标记值，其中，与所述每个样本问题向量对应的预设用户意图的标记值最大；

将所有所述样本问题向量分别输入第一双层循环神经网络进行编码，得到各个样本问题向量各自对应的语境特征向量；

将各个所述语境特征向量分别投入 Key-Value 记忆网络中进行迭代计算，直到达到预设的迭代次数或模型收敛，然后输出迭代计算后的各个语境特征向量，所述 Key-Value 记忆网络中参与迭代计算各个 Key 值为各个样本问题文本中主语和谓语对应的向量；

将迭代计算后的各个语境特征向量分别输入第二双层循环神经网络进行解码，得到各个样本结果向量；

针对每个样本结果向量，通过内容相似度计算网络计算所述每个样本结果向量与各个意图向量之间的相似度，得到所述每个样本结果向量对应的各个相似度值，作为样本概率值，所述各个意图向量是指所述各个预设用户意图向量化后的向量值；

以所述每个样本结果向量对应的各个样本概率值为调整目标，调整所述第一双层循环神经网络、所述 Key-Value 记忆网络、所述第二双层循环神经网络和所述内容相似度计算网络的网络参数，以最小化所述每个样本结果向量对应的各个样本概率值与各个目标标记值之间的误差，所述各个目标标记值是指所述每个样本结果向量对应的样本问题向量的各个标记值；

若每个样本结果向量对应的各个样本概率值与各个目标标记值之间的误差满足预设的训练终止条件，则确定所述深度学习模型已训练好。

13. 根据权利要求 12 所述的计算机设备，其特征在于，所述将所述目标问题向量作为输入投入至预先训练好的基于注意力的深度学习模型，得到所述深度学习模型输出的、与各个预设用户意图分别对应的各个问题概率值包括：

将所述目标问题向量输入所述第一双层循环神经网络进行编码，得到目标语境特征向量；

将所述目标语境特征向量投入所述 Key-Value 记忆网络中进行迭代计算，直到达到预设的迭代次数或模型收敛，然后输出迭代计算后的目标语境特征向量；

将迭代计算后的目标语境特征向量输入所述第二双层循环神经网络进行解码，得到目

标结果向量；

通过所述内容相似度计算网络计算所述目标结果向量与各个意图向量之间的相似度，得到所述目标结果向量对应的各个相似度值，作为各个问题概率值。

14. 根据权利要求 11 所述的计算机设备，其特征在于，所述对所述目标问题文本进行向量化处理，得到目标问题向量包括：

将所述目标问题文本中的各个目标字词分别转换为 GloVe 词向量，得到初始问题向量；

判断所述各个目标字词是否均被 GloVe 词向量覆盖；

若所述各个目标字词均被 GloVe 词向量覆盖，则确定所述初始问题向量为目标问题向量；

若所述各个目标字词中任一个目标字词未被 GloVe 词向量覆盖，则将未被 GloVe 词向量覆盖的目标字词转换为 TransE 词向量，得到补充向量；

将所述补充向量添加至所述初始问题向量，得到目标问题向量。

15. 根据权利要求 11 至 14 中任一项所述的计算机设备，其特征在于，在从所述各个预设用户意图中选取出问题概率值最大的预设用户意图，作为所述用户的真实问题意图之后，所述处理器执行所述计算机可读指令时还实现如下步骤：

获取预设的各个问题答案元组，每个问题答案元组由一个预设用户意图和与所述一个预设用户意图对应的答案组成，其中，所述各个问题答案元组的预设用户意图各不相同；

从所述各个问题答案元组中选取预设用户意图与所述真实问题意图相同的一个问题答案元组，作为命中元组；

将所述命中元组的答案反馈至所述用户。

16. 一个或多个存储有计算机可读指令的可读存储介质，其特征在于，所述计算机可读指令被一个或多个处理器执行时，使得所述一个或多个处理器执行如下步骤：

获取用户提问的目标问题语音；

对所述目标问题语音进行音转字处理，得到目标问题文本；

对所述目标问题文本进行向量化处理，得到目标问题向量；

将所述目标问题向量作为输入投入至预先训练好的基于注意力的深度学习模型，得到所述深度学习模型输出的、与各个预设用户意图分别对应的各个问题概率值，每个问题概率值各自表征了所述目标问题文本属于与所述每个问题概率值对应的预设用户意图的概率；

从所述各个预设用户意图中选取出问题概率值最大的预设用户意图，作为所述用户的真实问题意图。

17. 根据权利要求 16 所述的可读存储介质，其特征在于，所述深度学习模型包括第一双层循环神经网络、Key-Value 记忆网络、第二双层循环神经网络和内容相似度计算网络四个部分，所述深度学习模型通过以下步骤预先训练好：

收集属于所述各个预设用户意图的样本问题语音；

对收集到的样本问题语音分别进行音转字处理，得到样本问题文本；

对所述样本问题文本进行向量化处理，得到样本问题向量；

针对每个样本问题向量，为所述每个样本问题向量分别针对各个预设用户意图设定标记值，得到所述每个样本问题向量的各个标记值，其中，与所述每个样本问题向量对应的预设用户意图的标记值最大；

将所有所述样本问题向量分别输入第一双层循环神经网络进行编码，得到各个样本问题向量各自对应的语境特征向量；

将各个所述语境特征向量分别投入 Key-Value 记忆网络中进行迭代计算，直到达到预设的迭代次数或模型收敛，然后输出迭代计算后的各个语境特征向量，所述 Key-Value 记

忆网络中参与迭代计算各个 Key 值为各个样本问题文本中主语和谓语对应的向量；

将迭代计算后的各个语境特征向量分别输入第二双层循环神经网络进行解码，得到各个样本结果向量；

针对每个样本结果向量，通过内容相似度计算网络计算所述每个样本结果向量与各个意图向量之间的相似度，得到所述每个样本结果向量对应的各个相似度值，作为样本概率值，所述各个意图向量是指所述各个预设用户意图向量化后的向量值；

以所述每个样本结果向量对应的各个样本概率值为调整目标，调整所述第一双层循环神经网络、所述 Key-Value 记忆网络、所述第二双层循环神经网络和所述内容相似度计算网络的网络参数，以最小化所述每个样本结果向量对应的各个样本概率值与各个目标标记值之间的误差，所述各个目标标记值是指所述每个样本结果向量对应的样本问题向量的各个标记值；

若每个样本结果向量对应的各个样本概率值与各个目标标记值之间的误差满足预设的训练终止条件，则确定所述深度学习模型已训练好。

18. 根据权利要求 17 所述的可读存储介质，其特征在于，所述将所述目标问题向量作为输入投入至预先训练好的基于注意力的深度学习模型，得到所述深度学习模型输出的、与各个预设用户意图分别对应的各个问题概率值包括：

将所述目标问题向量输入所述第一双层循环神经网络进行编码，得到目标语境特征向量；

将所述目标语境特征向量投入所述 Key-Value 记忆网络中进行迭代计算，直到达到预设的迭代次数或模型收敛，然后输出迭代计算后的目标语境特征向量；

将迭代计算后的目标语境特征向量输入所述第二双层循环神经网络进行解码，得到目标结果向量；

通过所述内容相似度计算网络计算所述目标结果向量与各个意图向量之间的相似度，得到所述目标结果向量对应的各个相似度值，作为各个问题概率值。

19. 根据权利要求 16 所述的可读存储介质，其特征在于，所述对所述目标问题文本进行向量化处理，得到目标问题向量包括：

将所述目标问题文本中的各个目标字词分别转换为 GloVe 词向量，得到初始问题向量；

判断所述各个目标字词是否均被 GloVe 词向量覆盖；

若所述各个目标字词均被 GloVe 词向量覆盖，则确定所述初始问题向量为目标问题向量；

若所述各个目标字词中任一目标字词未被 GloVe 词向量覆盖，则将未被 GloVe 词向量覆盖的目标字词转换为 TransE 词向量，得到补充向量；

将所述补充向量添加至所述初始问题向量，得到目标问题向量。

20. 根据权利要求 16 至 19 中任一项所述的可读存储介质，其特征在于，在从所述各个预设用户意图中选取出问题概率值最大的预设用户意图，作为所述用户的真实问题意图之后，所述计算机可读指令被一个或多个处理器执行时，使得所述一个或多个处理器还执行如下步骤：

获取预设的各个问题答案元组，每个问题答案元组由一个预设用户意图和与所述一个预设用户意图对应的答案组成，其中，所述各个问题答案元组的预设用户意图各不相同；

从所述各个问题答案元组中选取预设用户意图与所述真实问题意图相同的一个问题答案元组，作为命中元组；

将所述命中元组的答案反馈至所述用户。

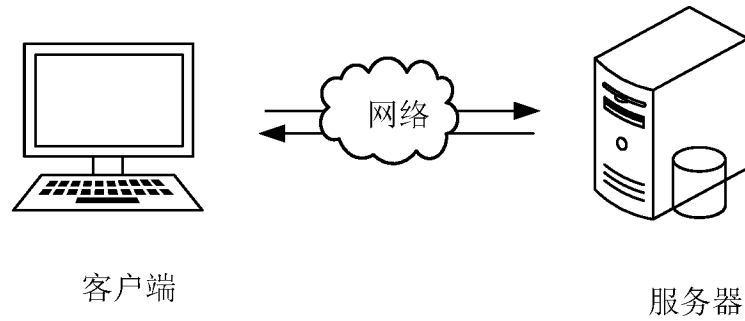


图 1

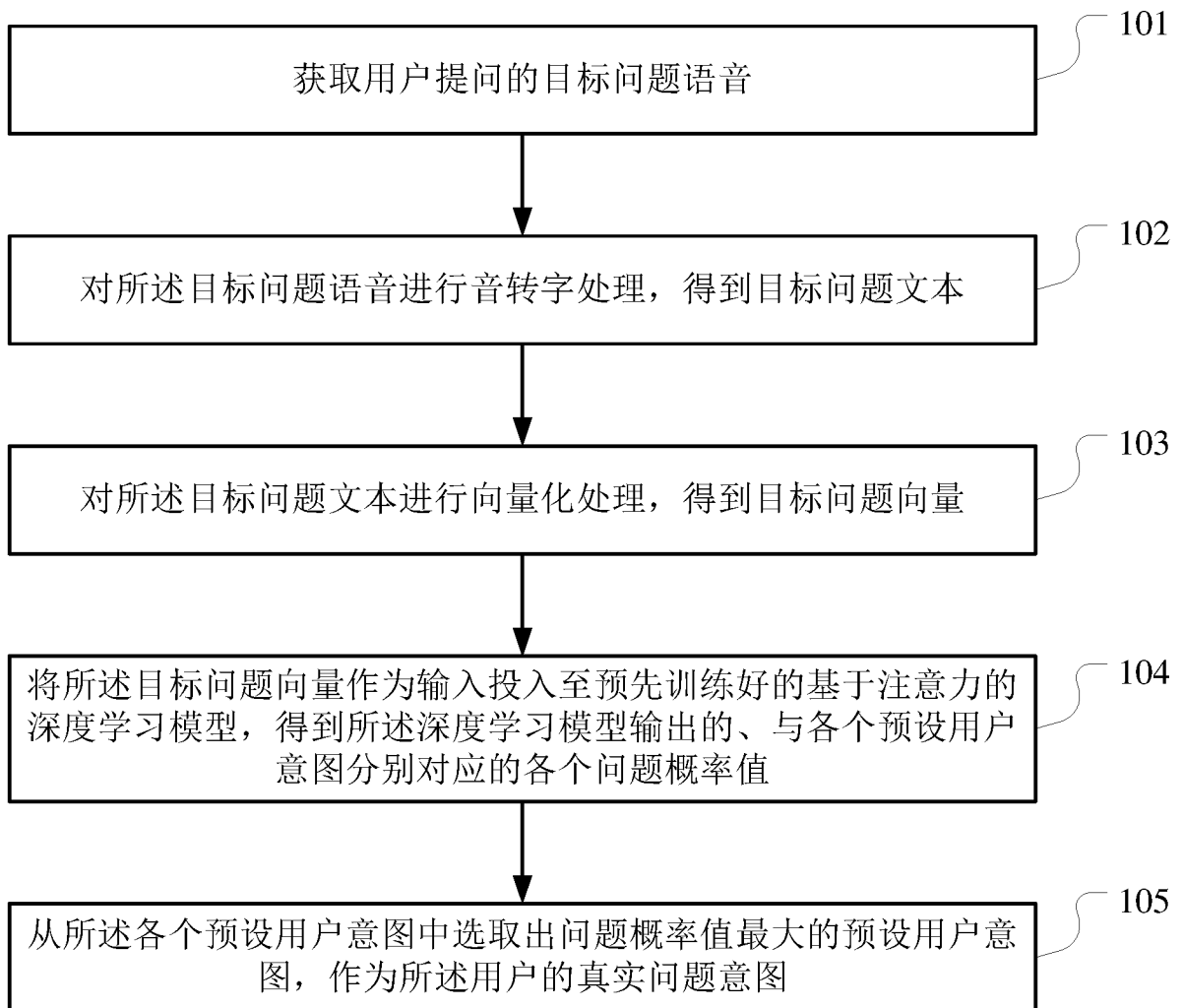


图 2

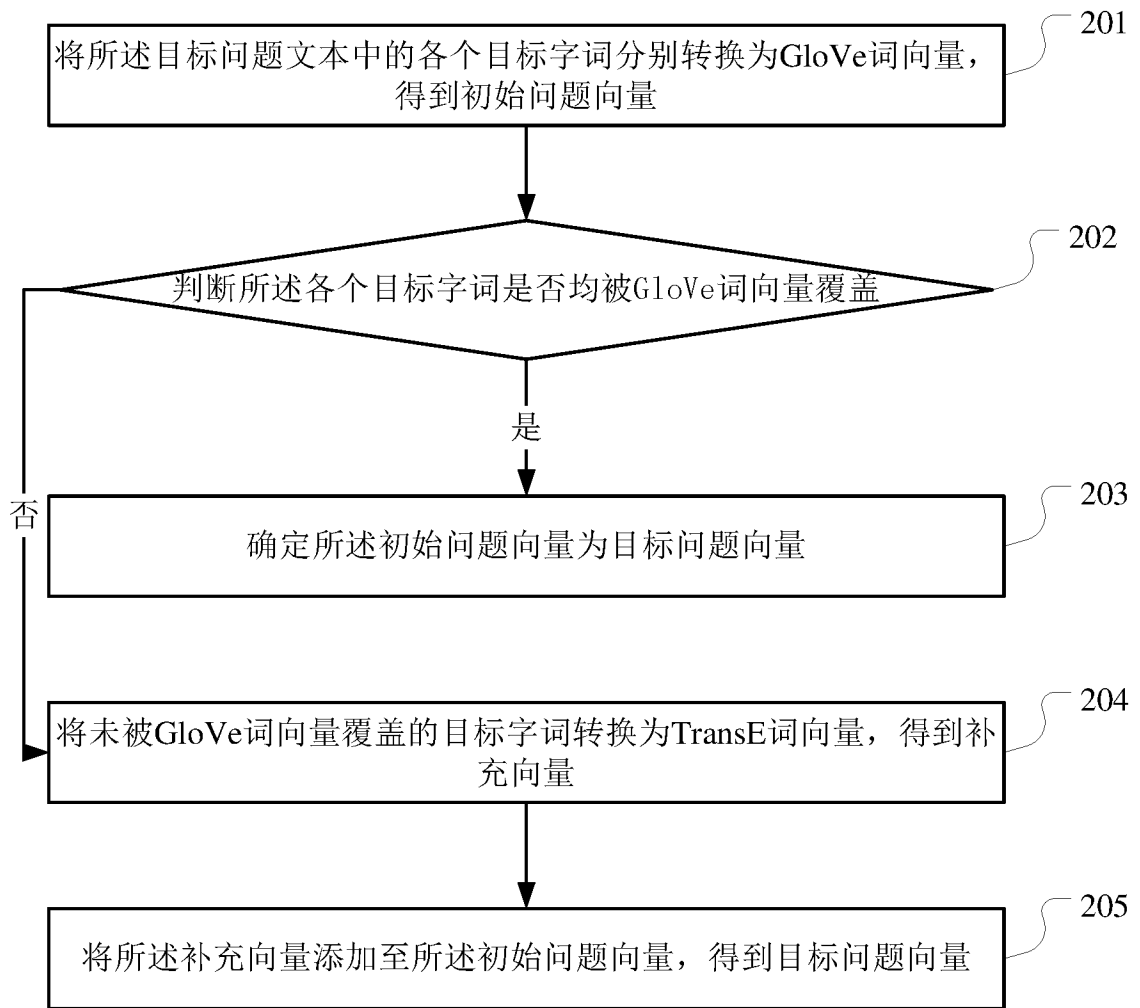


图 3

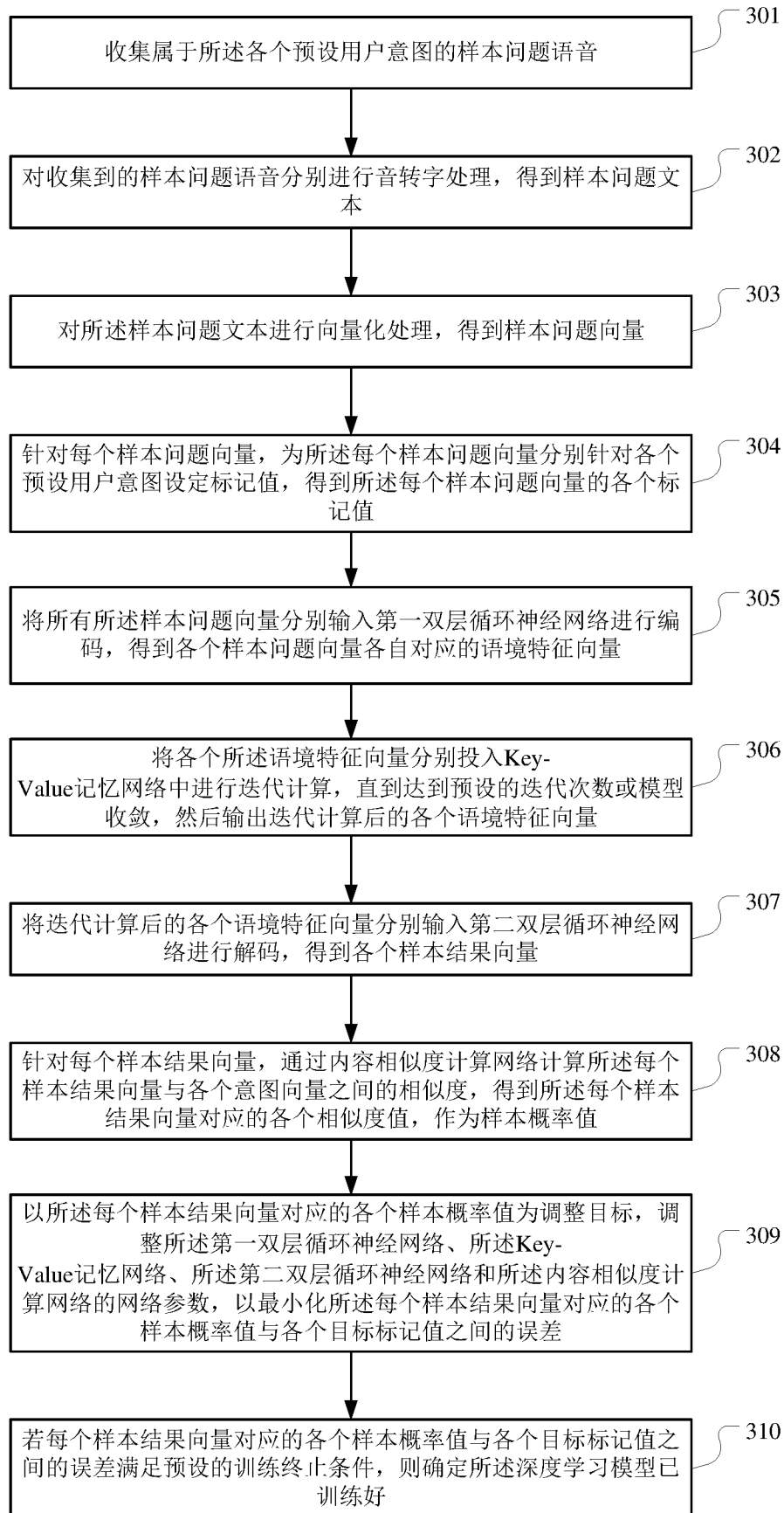


图 4

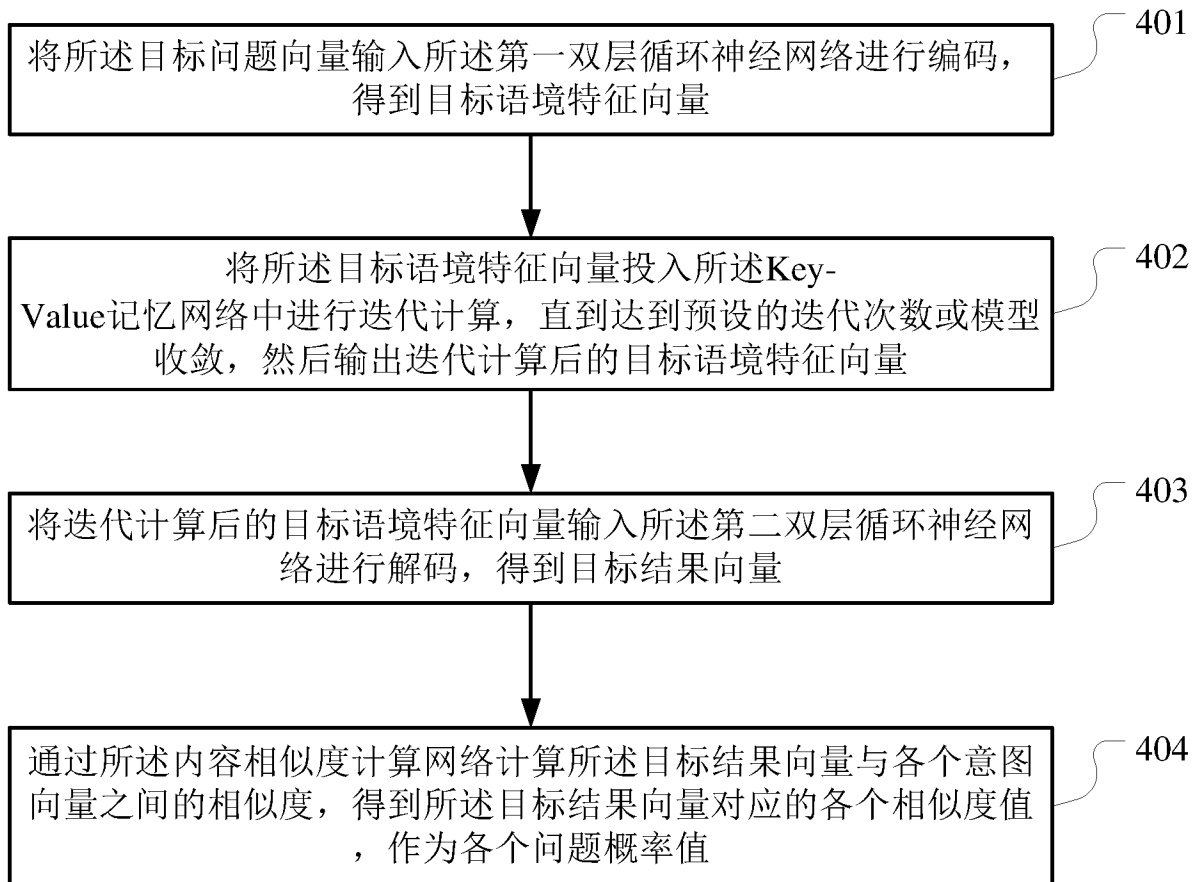


图 5

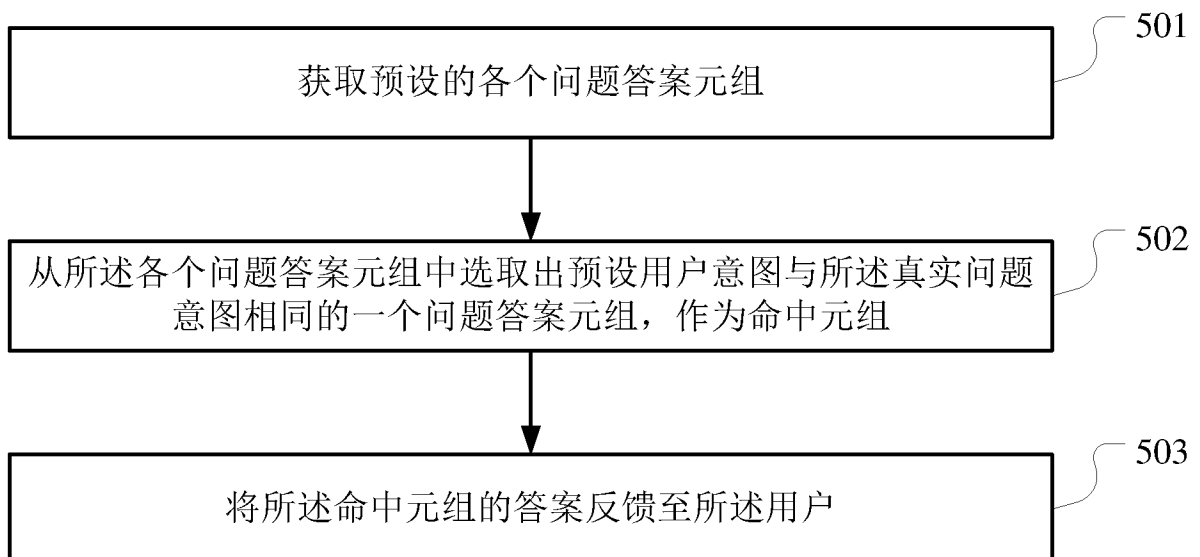


图 6

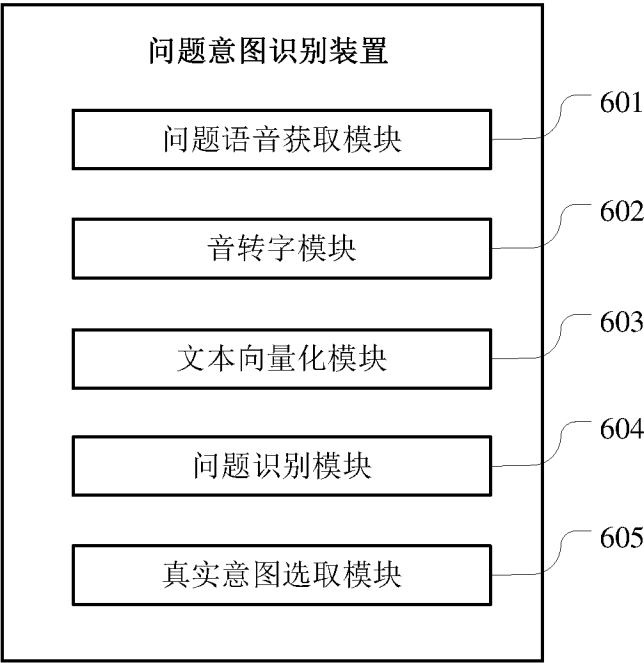


图 7

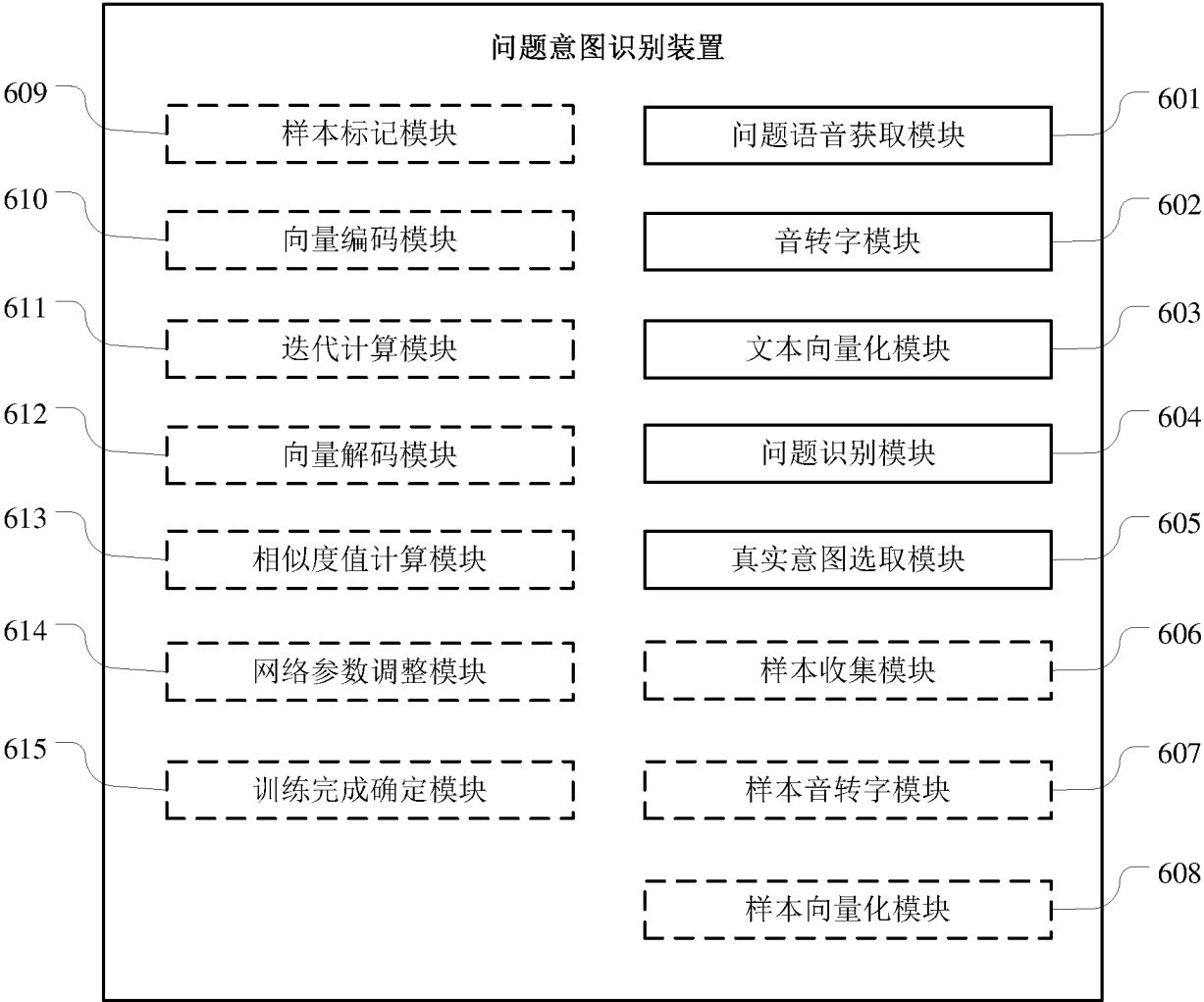


图 8

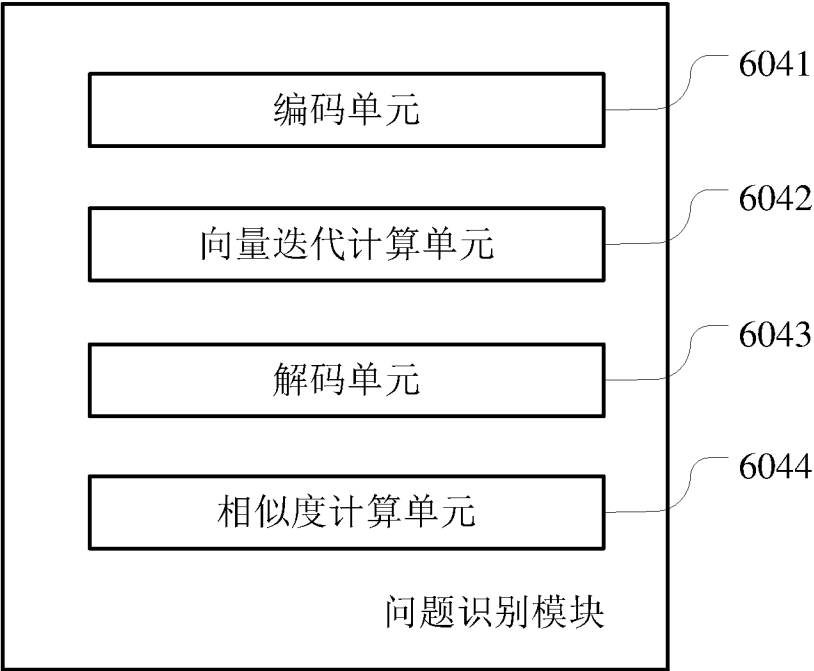


图 9

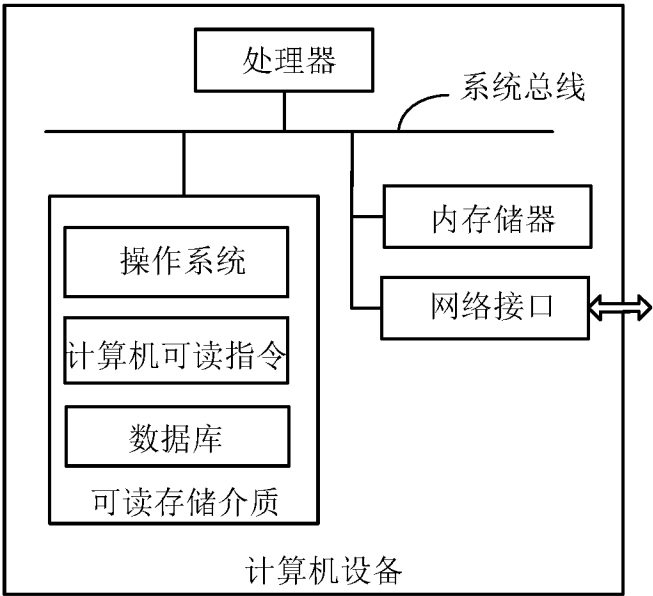


图 10

INTERNATIONAL SEARCH REPORT

International application No.

PCT/CN2019/102922

A. CLASSIFICATION OF SUBJECT MATTER

G06F 16/33(2019.01)i; G06F 16/332(2019.01)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)

G06F

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)

CNPAT; CNKI; WPI; EPODOC; IEEE: 语音, 识别, 文本, 向量, 注意力, 深度学习模型, 神经网络, 概率, 意图, 问题, 双层, voice, speech, recogni+, text, RNN, CNN, neural network, intent, vector?, attention, key-value

C. DOCUMENTS CONSIDERED TO BE RELEVANT

Category*	Citation of document, with indication, where appropriate, of the relevant passages	Relevant to claim No.
X	CN 108073574 A (SAMSUNG ELECTRONICS CO., LTD.) 25 May 2018 (2018-05-25) description, paragraphs [0051]-[0083], [0112]-[0134] and [0153]	1-20
X	US 2018144385 A1 (WAL-MART STORES, INC.) 24 May 2018 (2018-05-24) description, paragraphs [0047]-[0059]	1-20
Y	CN 109376361 A (BEIJING JIUHU TIMES INTELLIGENT TECHNOLOGY CO., LTD.) 22 February 2019 (2019-02-22) description, paragraphs [0046]-[0069]	1-20
Y	CN 109800306 A (SHENZHEN TCL NEW TECHNOLOGY CO., LTD.) 24 May 2019 (2019-05-24) description, paragraphs [0079]-[0137]	1-20
A	US 2019139537 A1 (KABUSHIKI KAISHA TOSHIBA) 09 May 2019 (2019-05-09) entire document	1-20
A	CN 108920622 A (BEIJING QIYI CENTURY SCIENCE AND TECHNOLOGY CO., LTD.) 30 November 2018 (2018-11-30) entire document	1-20



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

24 February 2020

Date of mailing of the international search report

03 March 2020

Name and mailing address of the ISA/CN

China National Intellectual Property Administration
No. 6, Xitucheng Road, Jimenqiao Haidian District, Beijing
100088
China

Authorized officer

Facsimile No. (86-10)62019451

Telephone No.

INTERNATIONAL SEARCH REPORT
Information on patent family members

International application No.

PCT/CN2019/102922

Patent document cited in search report			Publication date (day/month/year)	Patent family member(s)			Publication date (day/month/year)
CN	108073574	A	25 May 2018	KR	20180055189	A	25 May 2018
				US	2018137855	A1	17 May 2018
				EP	3324405	A1	23 May 2018
				JP	2018081298	A	24 May 2018
US	2018144385	A1	24 May 2018	None			
CN	109376361	A	22 February 2019	None			
CN	109800306	A	24 May 2019	None			
US	2019139537	A1	09 May 2019	JP	2019086679	A	06 June 2019
CN	108920622	A	30 November 2018	None			

国际检索报告

国际申请号

PCT/CN2019/102922

A. 主题的分类

G06F 16/33 (2019.01)i; G06F 16/332 (2019.01)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)

G06F

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))

CNPAT; CNKI; WPI; EPDOC; IEEE: 语音, 识别, 文本, 向量, 注意力, 深度学习模型, 神经网络, 概率, 意图, 问题, 双层, voice, speech, recogni+, text, RNN, CNN, neural network, intent, vector?, attention, key-value

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
X	CN 108073574 A (三星电子株式会社) 2018年 5月 25日 (2018 - 05 - 25) 说明书第[0051]-[0083]、[0112]-[0134]、[0153]段	1-20
X	US 2018144385 A1 (WAL-MART STORES, INC.) 2018年 5月 24日 (2018 - 05 - 24) 说明书第[0047]-[0059]段	1-20
Y	CN 109376361 A (北京九狐时代智能科技有限公司) 2019年 2月 22日 (2019 - 02 - 22) 说明书第[0046]-[0069]段	1-20
Y	CN 109800306 A (深圳TCL新技术有限公司) 2019年 5月 24日 (2019 - 05 - 24) 说明书第[0079]-[0137]段	1-20
A	US 2019139537 A1 (KABUSHIKI KAISHA TOSHIBA) 2019年 5月 9日 (2019 - 05 - 09) 全文	1-20
A	CN 108920622 A (北京奇艺世纪科技有限公司) 2018年 11月 30日 (2018 - 11 - 30) 全文	1-20

☐ 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 2月 24日

国际检索报告邮寄日期

2020年 3月 3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

向琳

电话号码 86-10-53961566

国际检索报告
关于同族专利的信息

国际申请号

PCT/CN2019/102922

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	108073574	A	2018年 5月 25日	KR	20180055189	A	2018年 5月 25日
				US	2018137855	A1	2018年 5月 17日
				EP	3324405	A1	2018年 5月 23日
				JP	2018081298	A	2018年 5月 24日
US	2018144385	A1	2018年 5月 24日	无			
CN	109376361	A	2019年 2月 22日	无			
CN	109800306	A	2019年 5月 24日	无			
US	2019139537	A1	2019年 5月 9日	JP	2019086679	A	2019年 6月 6日
CN	108920622	A	2018年 11月 30日	无			