

(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
2025 年 1 月 2 日 (02.01.2025)



(10) 国际公布号
WO 2025/000910 A1

- (51) 国际专利分类号:
G06F 16/35 (2019.01) **G06F 40/30** (2020.01)
G06F 40/289 (2020.01)
- (21) 国际申请号: PCT/CN2023/136792
- (22) 国际申请日: 2023 年 12 月 6 日 (06.12.2023)
- (25) 申请语言: 中文
- (26) 公布语言: 中文
- (30) 优先权:
202310800741.7 2023年6月30日 (30.06.2023) CN
- (71) 申请人: 中国电信股份有限公司(**CHINA TELECOM CORPORATION LIMITED**) [CN/CN]; 中国北京市西城区金融大街31号, Beijing 100033 (CN)。
- (72) 发明人: 付薇薇(**FU, Weiwei**); 中国北京市西城区金融大街31号, Beijing 100033 (CN)。刘欣璋(**LIU, Xinzhang**); 中国北京市西城区金融大街31号, Beijing 100033 (CN)。陈诣文(**CHEN, Yiwen**); 中国北京市西城区金融大街31号, Beijing 100033 (CN)。方瑞玉(**FANG, Ruiyu**); 中国北京市西城区金融大街31号, Beijing 100033

(CN)。宋双永(**SONG, Shuangyong**); 中国北京市西城区金融大街31号, Beijing 100033 (CN)。张寅(**ZHANG, Yin**); 中国北京市西城区金融大街31号, Beijing 100033 (CN)。

- (74) 代理人: 中国贸促会专利商标事务所有限有限公司(**CCPIT PATENT AND TRADEMARK LAW OFFICE**); 中国北京市西城区复兴门内大街158号远洋大厦F10层, Beijing 100031 (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, CV, CZ, DE, DJ, DK, DM, DO, DZ, EC, EE, EG, ES, FI, GB, GD, GE, GH, GM, GT, HN, HR, HU, ID, IL, IN, IQ, IR, IS, IT, JM, JO, JP, KE, KG, KH, KN, KP, KR, KW, KZ, LA, LC, LK, LR, LS, LU, LY, MA, MD, MG, MK, MN, MU, 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, CV, GH, GM, KE, LR, LS, MW, MZ,

(54) **Title:** KEYWORD CATEGORY IDENTIFICATION METHOD AND APPARATUS, AND ELECTRONIC DEVICE

(54) 发明名称: 识别关键词类别的方法、装置及电子设备

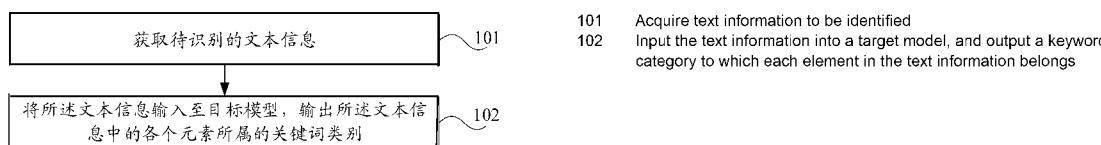


图 1

(57) **Abstract:** The present invention provides a keyword category identification method and apparatus, and an electronic device. The method comprises: acquiring text information to be identified; and inputting the text information into a target model, and outputting a keyword category to which each element in the text information belongs, wherein the target model comprises a first processing layer used for extracting semantic features of the text information, a second processing layer used for extracting, from the text information, a target keyword existing in a service dictionary, and a third processing layer used for determining, on the basis of the semantic features and the target keyword, the keyword category to which each element in the text information belongs, the first processing layer and the second processing layer are arranged in parallel, and the service dictionary comprises a plurality of keywords identified in a service process.

(57) 摘要: 本公开提供了一种识别关键词类别的方法、装置及电子设备。该方法包括: 获取待识别的文本信息; 将文本信息输入至目标模型, 输出文本信息中的各个元素所属的关键词类别; 其中, 目标模型包括用于提取文本信息的语义特征的第一处理层、用于提取文本信息中存在于业务词典中的目标关键词的第二处理层、以及用于基于语义特征和目标关键词确定文本信息中的各个元素所属的关键词类别的第三处理层, 第一处理层与第二处理层并行设置; 业务词典包括在业务过程中已识别出的多个关键词。

NA, RW, SC, 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, ME, 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))。

识别关键词类别的方法、装置及电子设备

相关申请的交叉引用

本申请是以 CN 申请号为 202310800741.7，申请日为 2023 年 6 月 30 日的申请为基础，并主张其优先权，该 CN 申请的公开内容在此作为整体引入本申请中。

技术领域

本公开涉及通信技术领域，尤其涉及一种识别关键词类别的方法、装置及电子设备。

背景技术

通信助理是移动互联网时代，为通信助理业务用户开发的一款增值服务的应用程序，提供漏话查询、通信录备份、短信备份、日程提醒、私密空间、黑白名单等功能。

目前，各个运营商的通信助理业务已日趋成熟，市场竞争也日趋激烈，因此，通信助理需要在业务上做到语义理解，同时帮助用户提取重要信息，并进行提醒。

其中，在通信助理业务中，从通话中提取关键信息至关重要。如果在从通话内容中提取关键词之前，能够确定该通话内容对应的文本信息中各个元素所属的类别，则可以更加准确地提取文本信息中的关键词。

发明内容

第一方面，本公开实施例提供了一种识别关键词类别的方法，所述方法包括：获取待识别的文本信息；和将所述文本信息输入至目标模型，输出所述文本信息中的各个元素所属的关键词类别；其中，所述目标模型包括用于提取所述文本信息的语义特征的第一处理层、用于提取所述文本信息中存在于业务词典中的目标关键词的第二处理层、以及用于基于所述语义特征和所述目标关键词确定所述文本信息中的各个元素所属的关键词类别的第三处理层，所述第一处理层与所述第二处理层并行设置；所述业务词典包括在业务过程中已识别出的多个关键词。

第二方面，本公开的实施例提供了一种识别关键词类别的装置，所述装置包括：第一获取模块，用于获取待识别的文本信息；和识别模块，用于将所述文本信息输入至目标模型，输出所述文本信息中的各个元素所属的关键词类别；其中，所述目标模

型包括用于提取所述文本信息的语义特征的第一处理层、用于提取所述文本信息中存在于业务词典中的目标关键词的第二处理层、以及用于基于所述语义特征和所述目标关键词确定所述文本信息中的各个元素所属的关键词类别的第三处理层，所述第一处理层与所述第二处理层并行设置；所述业务词典包括在业务过程中已识别出的多个关键词。

第三方面，本公开的实施例提供了一种电子设备，包括存储器、收发机和处理器：所述存储器用于存储计算机程序；所述收发机用于在所述处理器的控制下收发数据；所述处理器用于读取所述存储器中的计算机程序并执行上述第一方面所述的识别关键词类别的方法。

第四方面，本公开的实施例提供了一种处理器可读存储介质，所述处理器可读存储介质存储有计算机程序，所述计算机程序用于使处理器执行上述第一方面所述的识别关键词类别的方法。

第五方面，本公开的实施例提供了一种计算机程序，包括：指令，所述指令当由处理器执行时使所述处理器执行上述第一方面所述的识别关键词类别的方法。

附图说明

为了更清楚地说明本公开实施例的技术方案，下面将对本公开实施例的描述中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公开的一些实施例，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 为本公开实施例提供的识别关键词类别的方法的流程图；

图 2 为相关技术中进行预训练时的掩码（MASK）操作的示意图；

图 3 为本公开实施例中进行预训练时的 MASK 操作的示意图；

图 4 为本公开实施例的识别关键词类别的方法的具体实施方式中目标模型的架构示意图；

图 5 为本公开实施例提供的识别关键词类别的装置的结构框图；

图 6 为本公开实施例提供的电子设备的结构框图。

具体实施方式

本公开实施例中术语“和/或”，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：单独存在 A，同时存在 A 和 B，单独存在 B 这三种情况。字符“/”一般表示前后关联对象是一种“或”的关系。

本申请实施例中术语“多个”是指两个或两个以上，其它量词与之类似。

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，所描述的实施例仅仅是本申请一部分实施例，并不是全部的实施例。基于本申请中的实施例，本领域普通技术人员在没有做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。

本公开的发明人发现，在相关技术中，仅基于文本信息的语义特征来识别各个元素所属的类别，在通信助理的不同业务场景下可能会误识别，从而导致关键词的误提取，进而降低业务体验。

鉴于此，本公开的实施例提供了一种识别关键词类别的方法，以提升从文本信息中提取关键词的准确度。

第一方面，本公开的实施例提供了一种识别关键词类别的方法，如图 1 所示，该方法包括如下步骤 101 至 102。

步骤 101：获取待识别的文本信息。

在一些实施例中，所述文本信息可以为通话内容对应的文本信息，例如将通话内容转换为的文本内容；或者，所述文本信息也可以为从视频文件或者除通话内容之外的其他音频文件中提取的文本信息；或者所述文本信息也可以为直接获取的文本内容，例如文档中的内容。

步骤 102：将所述文本信息输入至目标模型，输出所述文本信息中的各个元素所属的关键词类别。

其中，所述目标模型包括用于提取所述文本信息的语义特征的第一处理层、用于提取所述文本信息中存在于业务词典中的目标关键词的第二处理层、以及用于基于所述语义特征和所述目标关键词确定所述文本信息中的各个元素所属的关键词类别的第三处理层，所述第一处理层与所述第二处理层并行设置；所述业务词典包括在业务过程中已识别出的多个关键词。

即文本信息输入至目标模型中，通过第一处理层获取该文本信息的语义特征，通过第二处理层从文本信息中提取存在于业务词典中的目标关键词，然后将获取的语义特征以及目标关键词，分别输入至第三处理层，从而通过第三处理层输出文本信息中各个元素所属的关键词的类别。

由此可知，在本公开的实施例中，可以将业务过程中已识别出的关键词存入业务词典，从而使得目标模型可以基于该业务词典从文本信息中提取关键词。这样，再进一步结合文本信息的语义特征，则可以更加准确地识别出文本信息中各个元素所属的关键词类别。

5 在一些实施例中，所述关键词类别包括时间关键词、地点关键词、人物关键词、事件关键词中的至少一种。可以理解的是，关键词类别还可以包括除此外列举的类别之外的其他类别。

另外，通过步骤 102 得到文本信息中各个元素所属的关键词类别之后，则可以根据各个元素所属的关键词类别，从文本信息中提取关键词，从而使得提取到的关键词更加
10 准确。

此外，上述业务词典可以由业务方给出，也可以在目标模型实际部署上线后，由数据回流收集得到的置信度和支持度均比较高的关键词构成。

由上述步骤 101 至 102 可知，在本公开实施例中，能够获取待识别的文本信息，从而将文本信息输入目标模型，输出文本信息中的各个元素所属的关键词类别，其中，所述目标模型包括用于提取所述文本信息的语义特征的第一处理层、用于提取所述文本信息中存在于业务词典中的目标关键词的第二处理层、以及用于基于所述语义特征和所述目标关键词确定所述文本信息中的各个元素所属的关键词类别的第三处理层，所述第一处理层与第二处理层并行设置；所述业务词典包括在业务过程中已识别出的多个关键词。
15

由此可见，在本公开实施例中，上述目标模型在用于提取文本信息的语义特征的第一处理层的基础上，增加了提取业务词典中的关键词的第二处理层，这样，目标模型则可以结合语义特征和业务过程中已经识别出的关键词，来确定文本信息中各个元素所属的关键词类别，充分利用了已有的关键词业务知识，从而能够提升确定文本信息中各个元素所属关键词类别的准确度，进而提升从文本信息中提取关键词的准确度。
20

25 在一些实施例中，所述第一处理层包括词表征层、段表征层、位置表征层以及多注意力网络层，所述词表征层、所述段表征层和所述位置表征层并行设置，且所述词表征层、所述段表征层和所述位置表征层的输出分别输入至所述多注意力网络层。

由此可知，文本信息输入至第一处理层后，分别通过词表征层、段表征层和位置表征层成进行处理，分别得到词向量、段向量和位置向量，进而词向量、段向量和位置向量
30 输入至多注意力网络层进行处理，输出文本语义向量。

这里，词表征（**Token Embedding**）层，将输入的文本信息的元素转换成固定维度的词向量，例如对于英文，将一个单词会被拆成词根和词缀，然后分别将词根和词缀转换成固定维度的向量；对于中文，将每一个字转换成固定维度的向量；这里转换得到的向量即为词向量。

5 段表征（**Segment Embedding**）层用于区别句子在文中的顺序，并转换成对应的向量，这里转换得到的向量即为段向量。

位置表征（**Position Embedding**）层用于将文本信息中元素的位置信息编码成特征向量，这里转换得到的向量即为位置向量。

多注意力网络（**Multi-attention**）层能够将多方面的特征进行融合，则在本公开的实
10 施例中，可以将通过词表征层、段表征层和位置表征层提取的这三方面的语义特征进行融合，进而得到能够更加准确地表征文本信息的语义特征的文本语义向量。

此外，上述第二处理层的维度与可以与 **Token Embedding** 层的维度相同，且初始化参数也与 **Token Embedding** 层相同；并且，还可以设置 **Keyword Embedding** 层的参数是可训练的。

15 在一些实施例中，所述方法还包括如下步骤 A1 至 A3。

步骤 A1：获取预训练文本信息。

步骤 A2：将所述预训练文本信息中存在于所述业务词典中的关键词替换为掩码（即进行 **MASK** 操作），得到预训练输入文本。

步骤 A3：根据所述预训练输入文本，对所述第一处理层进行预训练。

20 例如在上述目标模型使用之前，可以采用上述预训练输入文本，对第一处理层进行预训练。

例如预训练文本信息为“你好，我是 **AB** 外卖，请下来取您的快递”，且命中业务词典的关键词为“**AB** 外卖”，则可以将“**AB** 外卖”替换为掩码<**MASK**>，即得到：“你好，我是
25 <**MASK**><**MASK**><**MASK**><**MASK**>，请下来取您的快递”，则“你好，我是
<**MASK**><**MASK**><**MASK**><**MASK**>，请下来取您的快递”即为预训练输入文本。

其中，在相关技术中的 **MLM** 任务（即指定周围词来预测中心词的任务）中，基于双向编码表示器（**Bidirectional Encoder Representation Transformers, BERT**）的 **Multi-attention** 设置的自注意力机制（**self-attention**）层的 **MASK** 矩阵是随机生成的，如图 2 所示。即在相关技术中的 **BERT** 模型在做 **MLM** 预训练时，通过随机掩盖一些词（替换为
30 统一标记符[**MASK**]），然后预测这些被遮盖的词来训练双向语言模型，并且使每个词的

表征参考上下文信息，这样做会使得模型不能针对性的对关键词序列预测任务提升。

而本公开实施例中的预训练方法，对第一处理层的 MASK 矩阵做了特殊限制，仅仅对预训练输入文本中命中业务词典的关键词部分做 MASK 操作，例如图 3 所示，从而使得第一处理层可以针对关键词序列进行预测。

5 由此可知，相对于相关技术而言，在本公开实施例中的预训练过程中，进行 MASK 操作的内容具有针对性，则获得预训练输入文本所需的时间更短，这样，在同样的时间内，采用本公开实施例中的预训练方法，可以对更多的预训练文本信息进行处理，从而得到更多的预训练输入文本，进而进一步提升第一处理层的准确度。

10 可选地，步骤 A3“根据所述预训练输入文本，对所述第一处理层进行预训练”，包括如下步骤 B1 至 B3。

 步骤 B1：将所述预训练输入文本输入至所述第一处理层，输出对所述掩码的预测结果。

 步骤 B2：计算与同一预训练输入文本对应的预测结果与被替换的关键词之间的交叉熵。

15 步骤 B3：根据所述交叉熵，调整所述第一处理层的参数。

 其中，按照上述步骤 A1 至 A2 得到多个预训练输入文本之后，则可以将预训练文本输入至上述第一处理层，从而得到第一处理层的输出（即预训练输入文本中掩码被预测结果替换后内容），这样，则可以根据第一处理层的输出以及进行 MASK 操作之前的预训练文本信息，调整第一处理层的参数。例如计算二者的交叉熵，从而根据交叉熵来调整

20 第一处理层的参数。

 但是，这里第一处理层的输出（即预训练输入文本中掩码被预测结果替换后内容），与进行 MASK 操作之前的预训练文本信息存在重复的内容（例如前述示例：你好，我是，请下来取您的快递），这些重复的内容则无需再计算交叉熵，因此，为了简化计算过程，可以仅计算第一处理层输出的对掩码的预测结果，与被掩码替换的关键词之间的交叉熵。

25 另外，计算第一处理层输出的对掩码的预测结果，与被掩码替换的关键词之间的交叉熵，即为计算二者的损失函数值并进行梯度回传，从而使得第一处理层在掩码的位置的输出能预测出“被掩码替换的关键词”，例如上述示例中：“你好，我是 <MASK><MASK><MASK><MASK>，请下来取您的快递”，能够在“<MASK><MASK><MASK><MASK>”的位置处预测出“AB 外卖”。

30 可选地，上述步骤 102“将所述文本信息输入至目标模型，输出所述文本信息中的各

个元素所属的关键词类别”，包括如下步骤 C1 至 C3。

步骤 C1：将所述文本信息输入至所述第一处理层，输出用于表征所述文本信息的语义特征的文本语义向量。

5 步骤 C2：将所述文本信息输入至所述第二处理层，通过所述第二处理层提取所述目标关键词，将所述目标关键词拼接为目标长度得到目标信息，并将所述目标信息转换为嵌入（Embedding）向量，其中，所述目标长度为所述文本信息的长度。

步骤 C3：将所述文本语义向量和所述嵌入向量分别输入至所述第三处理层，输出所述文本信息中的各个元素所属的关键词类别。

其中，在步骤 C2 中，例如文本信息为“你好，我是 AB 外卖，请下来取您的快递。”，
10 且其中命中业务词典的目标关键词为“AB 外卖”，则可以将该目标关键词填充至与文本信息相同的长度，即：AB 外卖 [PAD] [PAD] [PAD] [PAD] [PAD] [PAD] [PAD] [PAD] [PAD] [PAD] [PAD] [PAD] [PAD] [PAD]；然后将填充后得到的目标信息转换为 Embedding 向量。

由此可知，通过步骤 C1 至 C3，将文本信息中命中业务词典的目标关键词与表征文本信息的语义特征的文本语义向量相结合，从而能够更加准确的得到文本信息中各个元素所属的关键词类别。

在一些实施例中，所述第三处理层包括至少一个标准全连接层和归一化处理层；

步骤 C3“将所述文本语义向量和所述嵌入向量分别输入至所述第三处理层，输出所述文本信息中的各个元素所属的关键词类别”，包括如下步骤 D1 至 D4。

20 步骤 D1：将所述文本语义向量与所述嵌入向量相加，得到第一向量；

步骤 D2：将所述第一向量输入至所述至少一个标准全连接层，输出第二向量；

步骤 D3：将所述第二向量输入至所述归一化处理层，输出所述文本信息中各个元素属于预先确定的多个关键词类别的概率（probs）分布；

步骤 D4：根据所述概率分布，输出所述文本信息中的各个元素所属的关键词类别。

25 其中，在步骤 D4 中，可以对文本信息中各个元素属于预先确定的多个关键词类别的 probs 分布取 argmax 操作（即在概率分布中，确定出各个元素对应的最大概率值），从而可以得到文本信息中的各个元素所属的关键词类别。

这里，argmax 是一种函数，是对函数求参数(集合)的函数。当存在一个函数 $y=f(x)$ 时，若有结果 $x0=\text{argmax}(f(x))$ ，则表示当函数 $f(x)$ 取 $x=x0$ 的时候，得到 $f(x)$ 取值范围的最大值；若有多个点使得 $f(x)$ 取得相同的最大值，那么 $\text{argmax}(f(x))$ 的结果就是一个点集。

即 $\text{argmax}(f(x))$ 是使得 $f(x)$ 取得最大值所对应的变量点 x (或 x 的集合)。因此, 在本公开实施例中, 对文本信息中各个元素属于预先确定的多个关键词类别的 **probs** 分布取 **argmax** 操作, 即为确定出各个元素对应的最大概率值。

另外, 标准全连接层在这里起到分类器的作用, 即文本语义向量和嵌入向量相加后的第一向量输入至标准全连接层后, 通过标准全连接层的处理, 将文本信息的各个元素划分为多个关键词类别, 进而通过归一化层的处理, 则可以得到文本信息中各个元素属于预先确定的多个关键词类别的概率分布, 这样, 在概率分布中, 各个元素对应的最大概率值所属的关键词类别, 即为该元素所属的关键词类别。

此外, 所述至少一个标准全连接层例如可以包括两个第一标准全连接层和一个第二标准全连接层, 所述第一标准全连接层的输出维度 (**outdim**) 设置可以与前文所述的词表征层中的隐藏层尺寸 (**hidden size**) 相同, 其余参数可默认; 所述第二标准全连接层的输入维度可以与前文所述的词表征层中的隐藏层尺寸 (**hidden size**) 相同, 输出维度为 $k*2+1$, k 表示预先确定的关键词类别的数量。这里, “ $k*2+1$ ” 中 “ $*2$ ” 是因为采用了 **BIO** 结构, 每种类别有 **begin** 和 **I** 的区别; “ $+1$ ” 是因为非关键词的文本标注为 **O**。

在一些实施例中, 所述将所述文本信息输入至目标模型之前, 所述方法还包括如下步骤 **E1** 至 **E2**。

步骤 **E1**: 获取所述文本信息所属的领域信息。

步骤 **E2**: 将所述领域信息添加至所述文本信息中。

其中, 可以采用相关技术中的任一方法获取文本信息所属的领域信息。

另外, 可以将领域信息添加至文本信息的前端, 二者可以采用预设字符 (例如句号 “。”) 连接。

此外, 相关技术中的序列标注技术在通讯助理这种领域细分程度高, 关键词的槽位抽取涉及多种不同领域, 类别数量大, 长尾数据分布严重的场景上, 容易因为不同场景内的槽位类型语义相似度较大, 在领域间误判。而本公开的实施例中, 可以在文本信息中融入领域信息, 从而可以在识别文本信息中的元素所属的关键词类别的情况下, 可以结合该领域信息, 进而降低误判几率, 进而提升关键词的提取准确度。

可以理解的是, 在前文所述的预训练过程中, 也可以融入领域信息, 例如在步骤 **A1** 之后, 且步骤 **A2** 之前, 获取预训练文本信息的领域信息, 并添加至该预训练文本信息中。

在一些实施例中, 所述方法还包括如下步骤 **F1** 至 **F2**。

步骤 **F1**: 采用序列标注方法, 对训练文本信息中的元素进行标注, 得到每一个所述

训练文本信息的标签。

步骤 F2：根据所述训练文本信息以及所述训练文本信息的标签，对所述目标模型进行训练。

需要说明的是，这里对目标模型的训练过程，与前文所述的预训练的过程是不同的两个过程。即预训练是对第一处理层进行的，在对第一处理层预训练完成之后，在第一处理层的基础上添加上述第二处理层和第三处理层，进而构建得到目标模型之后，再对目标模型进行训练。

其中，序列标注方法即为：给定一个序列，对序列中的每个元素做一个标记，或者说给每一个元素打一个标签。打标使用标准的序列标注 BIO 结构(即 B-begin, I-inside, O-outside)，例如训练文本信息为“你好，我是 AB 外卖，请下来取您的快递。”，标签则为：OOOOOB-外卖名称 I-外卖名称 I-外卖名称 I-外卖名称 OOOOOOOOOO。

此外，上述训练文本信息中也可以添加该训练文本信息的领域信息，例如上述训练文本信息为“你好，我是 AB 外卖，请下来取您的快递。”，添加领域信息“外卖”之后，对应的标签则为：OOOOOOOOOB-外卖名称 I-外卖名称 I-外卖名称 I-外卖名称 OOOOOOOOOO。

可见，在目标模型的训练阶段，一个训练样本包括训练文本信息、训练文本信息的领域信息和标签。

其中，在上述步骤 F2 中，“根据所述训练文本信息以及所述训练文本信息的标签，对所述目标模型进行训练”，具体可以为：

将训练文本信息输入至目标模型，得到对应的输出，从而对比模型的输出与训练文本对应的标签，可以确定目标模型预测的训练文本信息中各个元素所属的关键词类别是否准确，进而可以基于此来调整目标模型的参数直到目标模型输出的结果的准确度可以达到一定要求为止，完成目标模型的训练。

综上所述，本公开实施例的识别关键词类别的方法的具体实施方式可如下所述：

(1) 预训练阶段：

使用基础 BERT 中间层架构（即 Token Embedding 层、Segment Embedding 层、Position Embedding 层和 Multi-attention 层），作为第一处理层进行预训练，具体过程如下步骤 1.1 至 1.5 所述。

1.1 在其中一个预训练文本信息中融入领域信息，例如为：“外卖。你好，我是 AB 外卖，请下来取您的快递。”。

1.2 识别融入领域信息后的预训练文本信息中，命中业务词典的关键词的位置，例如“外卖。你好，我是 AB 外卖，请下来取您的快递。”，命中业务词典的关键词为 AB 外卖，索引位置为：[8,12]。

1.3 将命中业务词典的关键词的位置变为<MASK>，得到预训练输入文本，即得到：
5 “外卖。你好，我是<MASK><MASK><MASK><MASK>，请下来取您的快递。”。

1.4 将预训练输入文本输入至第一处理层（即分别输入至 Token Embedding 层、Segment Embedding 层、Position Embedding 层，然后这三层的输出再输入至 Multi-attention 层），输出对<MASK><MASK><MASK><MASK>的预测结果 Y(y1,y2,y3,y4)。

1.5 计算 Y (y1,y2,y3,y4) 与原始文本即 (A,B,外,卖) 的损失函数值并进行梯度回传，
10 使得第一处理层在<MASK><MASK><MASK><MASK>的输出能预测出“A,B,外,卖”四个字。

1.6 所有预训练文本信息上重复上述 1.1 至 1.5 的过程，从而完成预训练过程。

(2) 目标模型构建及训练阶段：

即对第一处理层进行预训练完成后，在第一处理层（即 Token Embedding 层、
15 Segment Embedding 层、Position Embedding 层和 Multi-attention 层）的基础上，进一步添加 Keyword Embedding 层（即第二处理层）、两个第一标准全连接层、一个第二标准全连接层、归一化处理层和结果输出层，从而得到如图 4 所示的目标模型。

构建完成目标模型之后，则对构建的目标模型进行训练，训练过程如下步骤 2.1 至 2.3 所述。

20 2.1 获取训练文本信息，并在训练文本信息中融入领域信息，例如为：“外卖。你好，我是 AB 外卖，请下来取您的快递。”。

2.2 采用序列标注方法，对训练文本信息中的元素进行标注，得到每一个所述训练文本信息的标签，例如“外卖。你好，我是 AB 外卖，请下来取您的快递。”，得到的标签为：
O O O O O B-外卖名称 I-外卖名称 I-外卖名称 I-外卖名称 O O O O O O O O O O。

25 2.3 将训练文本信息输入至目标模型，得到对应的输出，从而对比模型的输出与训练文本对应的标签，可以确定目标模型预测的训练文本信息中各个元素所属的关键词类别是否准确，进而可以基于此来调整目标模型的参数直到目标模型输出的结果的准确度可以达到一定要求为止，完成目标模型的训练。

(3) 目标模型的使用阶段：

30 获取到待识别的文本信息之后，通过如下步骤 3.1 至 3.5 的过程，得到该文本信息中

各个元素所属的关键词类别。

31.在文本信息中融入领域信息，例如文本信息为“你好，我是 AB 外卖，请下来取您的快递。”，融入领域信息后为：“外卖。你好，我是 AB 外卖，请下来取您的快递。”。

3.1：如图 4 所示，将融入领域信息后的文本信息分别输入至目标模型的 **Token Embedding** 层、**Segment Embedding** 层、**Position Embedding** 层、**Keyword Embedding** 层，分别得到词向量、段向量、位置向量，然后这三个向量分别输入至 **Multi-attention** 层，输出文本语义向量。

其中，**Keyword Embedding** 层从文本信息中提取命中业务词典的关键词，然后拼接至与文本信息相同的长度，再将其转换为嵌入向量。

3.2 文本语义向量和嵌入向量相加后，得到第一向量。

3.3 第一向量输入顺序输入至两个第一标准全连接层和一个第二标准全连接层，输出第二向量。

3.4 第二向量输入至归一化处理层，输出文本信息中各个元素属于预先确定的多个关键词类别的概率分布。

3.5.结果输出层，根据所述概率分布，输出所述文本信息中的各个元素所属的关键词类别。

综上所述，本公开的实施例具有如下优点：

1.将待识别的文本信息的领域信息融入至该文本信息中，从而一同输入至目标模型，从而便于在目标模型中，对细分领域实现知识引导，尽可能地解决领域间相似实体误判的问题。

2.在目标模型中，输入的每一个 **Token**（即对词的嵌入），它的表征由其对应的词向量、段向量和位置向量以及命中业务词典的嵌入向量相加产生。并且，给出了 **Keyword Embedding** 与 **Multi-attention** 层输出的向量融合的计算方法，实现了 **Keyword Embedding** 层与上述第一处理层的结合。

其中，由于业务特殊性，关键词在通话中有很多重复性，但是在相关技术中，这些已经整理识别的关键词语义没有有效利用途径。

由上述可知，为了解决这个问题，在本公开的实施例中，在输入层设置了 **Keyword Embedding** 层，能够充分利用已有的关键词业务知识，这样，在训练和后续运营中都能使已有知识得到最大发挥。

3.本公开的实施例中，对上述第一处理层进行预训练的过程中，仅对命中业务词典的

关键词进行 MAS 操作，这样可对大量数据进行特殊关键词的语义知识嵌入训练，给这些关键词赋予特殊的 **embedding** 参数值，帮助提升模型准确度。即本公开的实施例中，通过独特的 MASK 方式，能将某些确定的关键词信息融入到对话语料中，从而可以提升模型泛化性，提升模型预测准确率。

5 以上介绍了本公开实施例提供的识别关键词类别的方法，下面将结合附图介绍本公开实施例提供的识别关键词类别的装置。

参见图 5，本公开实施例还提供了一种识别关键词类别的装置，如图 5 所示，所述装置包括：第一获取模块 501 和识别模块 502。

第一获取模块 501，用于获取待识别的文本信息。

10 识别模块 502，用于将所述文本信息输入至目标模型，输出所述文本信息中的各个元素所属的关键词类别；

其中，所述目标模型包括用于提取所述文本信息的语义特征的第一处理层、用于提取所述文本信息中存在于业务词典中的目标关键词的第二处理层、以及用于基于所述语义特征和所述目标关键词确定所述文本信息中的各个元素所属的关键词类别的第三处理层，
15 所述第一处理层与所述第二处理层并行设置。

所述业务词典包括在业务过程中已识别出的多个关键词。

在一些实施例中，所述第一处理层包括词表征层、段表征层、位置表征层以及多注意力网络层，所述词表征层、所述段表征层和所述位置表征层并行设置，且所述词表征层、所述段表征层和所述位置表征层的输出分别输入至所述多注意力网络层。

20 在一些实施例中，所述装置还包括：第二获取模块，用于获取预训练文本信息；掩码处理模块，用于将所述预训练文本信息中存在于所述业务词典中的关键词替换为掩码，得到预训练输入文本；和预训练模块，用于根据所述预训练输入文本，对所述第一处理层进行预训练。

在一些实施例中，所述预训练模块具体用于：将所述预训练输入文本输入至所述第一处理层，输出对所述掩码的预测结果；计算与同一预训练输入文本对应的预测结果与被替换的关键词之间的交叉熵；和根据所述交叉熵，调整所述第一处理层的参数。

25 在一些实施例中，所述识别模块 502 包括：第一处理子模块，用于将所述文本信息输入至所述第一处理层，输出用于表征所述文本信息的语义特征的文本语义向量；第二处理子模块，用于将所述文本信息输入至所述第二处理层，通过所述第二处理层提取所述
30 目标关键词，将所述目标关键词拼接为目标长度得到目标信息，并将所述目标信息转换

为嵌入向量，其中，所述目标长度为所述文本信息的长度；和第三处理子模块，用于将所述文本语义向量和所述嵌入向量分别输入至所述第三处理层，输出所述文本信息中的各个元素所属的关键词类别。

在一些实施例中，所述第三处理层包括至少一个标准全连接层和归一化处理层；所述第三处理子模块具体用于：将所述文本语义向量与所述嵌入向量相加，得到第一向量；将所述第一向量输入至所述至少一个标准全连接层，输出第二向量；将所述第二向量输入至所述归一化处理层，输出所述文本信息中各个元素属于预先确定的多个关键词类别的概率分布；和根据所述概率分布，输出所述文本信息中的各个元素所属的关键词类别。

在一些实施例中，所述装置还包括：第三获取模块，用于获取所述文本信息所属的领域信息；和领域融合模块，用于将所述领域信息添加至所述文本信息中。

在一些实施例中，所述装置还包括：标注模块，用于采用序列标注方法，对训练文本信息中的元素进行标注，得到每一个所述训练文本信息的标签；和

训练模块，用于根据所述训练文本信息以及所述训练文本信息的标签，对所述目标模型进行训练。

在一些实施例中，所述关键词类别包括时间关键词、地点关键词、人物关键词、事件关键词中的至少一种。

其中，方法和装置是基于同一申请构思的，由于方法和装置解决问题的原理相似，因此装置和方法的实施可以相互参见，重复之处不再赘述。

需要说明的是，本申请实施例中对单元的划分是示意性的，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。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。

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，可以存储在一个处理器可读取存储介质中。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）或处理器（processor）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，ROM）、随机存取存储器（Random Access Memory，

RAM)、磁碟或者光盘等各种可以存储程序代码的介质。

在此需要说明的是,本公开实施例提供的上述装置,能够实现上述方法实施例所实现的所有方法步骤,且能够达到相同的技术效果,在此不再对本实施例中与方法实施例相同的部分及有益效果进行具体赘述。

5 本公开的实施例还提供了一种电子设备,如图6所示,该电子设备包括存储器620、收发机610、处理器600。

存储器620,用于存储计算机程序。

收发机610,用于在处理器600的控制下接收和发送数据。

10 处理器600用于读取所述存储器620中的计算机程序并执行前述第一方面所述的识别关键词类别的方法。

其中,在图6中,总线架构可以包括任意数量的互联的总线和桥,具体由处理器600代表的一个或多个处理器和存储器620代表的存储器的各种电路链接在一起。总线架构还可以将诸如外围设备、稳压器和功率管理电路等之类的各种其他电路链接在一起,这些都是本领域所公知的,因此,本文不再对其进行进一步描述。总线接口提供接口。收发15 机610可以是多个元件,即包括发送机和接收机,提供用于在传输介质上与各种其他装置通信的单元,这些传输介质包括无线信道、有线信道、光缆等传输介质。处理器600负责管理总线架构和通常的处理,存储器620可以存储处理器600在执行操作时所使用的数据。

处理器600可以是中央处理器(CPU)、专用集成电路(Application Specific Integrated20 Circuit, ASIC)、现场可编程门阵列(Field-Programmable Gate Array, FPGA)或复杂可编程逻辑器件(Complex Programmable Logic Device, CPLD),处理器600也可以采用多核架构。

在此需要说明的是,本公开实施例提供的上述装置,能够实现上述方法实施例所实现的所有方法步骤,且能够达到相同的技术效果,在此不再对本实施例中与方法实施例25 相同的部分及有益效果进行具体赘述。

在本公开的一些实施例中,还提供了一种计算机程序,包括:指令,所述指令当由处理器执行时使所述处理器执行前面所述的识别关键词类别的方法。

本公开的实施例还提供了一种处理器可读存储介质,所述处理器可读存储介质存储有计算机程序,所述计算机程序用于使所述处理器执行上述所述的识别关键词类别的方30 法。例如,该处理器可读存储介质为非瞬时性处理器可读存储介质。

所述处理器可读存储介质可以是处理器能够存取的任何可用介质或数据存储设备，包括但不限于磁性存储器（例如软盘、硬盘、磁带、磁光盘（MO）等）、光学存储器（例如 CD、DVD、BD、HVD 等）、以及半导体存储器（例如 ROM、EPROM、EEPROM、非易失性存储器（NAND FLASH）、固态硬盘（SSD））等。

5 本领域内的技术人员应明白，本申请的实施例可提供为方法、系统、或计算机程序产品。因此，本申请可采用完全硬件实施例、完全软件实施例、或结合软件和硬件方面的实施例的形式。而且，本申请可采用在一个或多个其中包含有计算机可用程序代码的计算机可用存储介质（包括但不限于磁盘存储器和光学存储器等）上实施的计算机程序产品的形式。

10 本申请是参照根据本申请实施例的方法、设备（系统）、和计算机程序产品的流程图和 / 或方框图来描述的。应理解可由计算机可执行指令实现流程图和 / 或方框图中的每一流程和 / 或方框、以及流程图和 / 或方框图中的流程和 / 或方框的结合。可提供这些计算机可执行指令到通用计算机、专用计算机、嵌入式处理机或其他可编程数据处理设备的处理器以产生一个机器，使得通过计算机或其他可编程数据处理设备的处理器执行的指令产生用于实现在流程图一个流程或多个流程和 / 或方框图一个方框或多个方框中指定的功能的装置。

15 这些处理器可执行指令也可存储在能引导计算机或其他可编程数据处理设备以特定方式工作的处理器可读存储器中，使得存储在该处理器可读存储器中的指令产生包括指令装置的制造品，该指令装置实现在流程图一个流程或多个流程和 / 或方框图一个方框或多个方框中指定的功能。

20 这些处理器可执行指令也可装载到计算机或其他可编程数据处理设备上，使得在计算机或其他可编程设备上执行一系列操作步骤以产生计算机实现的处理，从而在计算机或其他可编程设备上执行的指令提供用于实现在流程图一个流程或多个流程和 / 或方框图一个方框或多个方框中指定的功能的步骤。

25 显然，本领域的技术人员可以对本申请进行各种改动和变型而不脱离本申请的精神和范围。这样，倘若本申请的这些修改和变型属于本申请权利要求及其等同技术的范围之内，则本申请也意图包含这些改动和变型在内。

权 利 要 求

1.一种识别关键词类别的方法，包括：

获取待识别的文本信息；和

将所述文本信息输入至目标模型，输出所述文本信息中的各个元素所属的关键词类别；

其中，所述目标模型包括用于提取所述文本信息的语义特征的第一处理层、用于提取所述文本信息中存在于业务词典中的目标关键词的第二处理层、以及用于基于所述语义特征和所述目标关键词确定所述文本信息中的各个元素所属的关键词类别的第三处理层，所述第一处理层与所述第二处理层并行设置；

所述业务词典包括在业务过程中已识别出的多个关键词。

2.根据权利要求1所述的方法，其中，所述第一处理层包括词表征层、段表征层、位置表征层以及多注意力网络层，所述词表征层、所述段表征层和所述位置表征层并行设置，且所述词表征层、所述段表征层和所述位置表征层的输出分别输入至所述多注意力网络层。

3.根据权利要求1或2所述的方法，还包括：

获取预训练文本信息；

将所述预训练文本信息中存在于所述业务词典中的关键词替换为掩码，得到预训练输入文本；和

根据所述预训练输入文本，对所述第一处理层进行预训练。

4.根据权利要求3所述的方法，其中，所述根据所述预训练输入文本，对所述第一处理层进行预训练，包括：

将所述预训练输入文本输入至所述第一处理层，输出对所述掩码的预测结果；

计算与同一预训练输入文本对应的预测结果与被替换的关键词之间的交叉熵；和
根据所述交叉熵，调整所述第一处理层的参数。

5.根据权利要求1至4任意一项所述的方法，其中，所述将所述文本信息输入至

目标模型，输出所述文本信息中的各个元素所属的关键词类别，包括：

将所述文本信息输入至所述第一处理层，输出用于表征所述文本信息的语义特征的文本语义向量；

将所述文本信息输入至所述第二处理层，通过所述第二处理层提取所述目标关键词，将所述目标关键词拼接为目标长度得到目标信息，并将所述目标信息转换为嵌入向量，其中，所述目标长度为所述文本信息的长度；和

将所述文本语义向量和所述嵌入向量分别输入至所述第三处理层，输出所述文本信息中的各个元素所属的关键词类别。

6.根据权利要求 5 所述的方法，其中，所述第三处理层包括至少一个标准全连接层和归一化处理层；

所述将所述文本语义向量和所述嵌入向量分别输入至所述第三处理层，输出所述文本信息中的各个元素所属的关键词类别，包括：

将所述文本语义向量与所述嵌入向量相加，得到第一向量；

将所述第一向量输入至所述至少一个标准全连接层，输出第二向量；

将所述第二向量输入至所述归一化处理层，输出所述文本信息中各个元素属于预先确定的多个关键词类别的概率分布；和

根据所述概率分布，输出所述文本信息中的各个元素所属的关键词类别。

7.根据权利要求 1 至 6 任意一项所述的方法，还包括：

在将所述文本信息输入至目标模型之前，获取所述文本信息所属的领域信息；和
将所述领域信息添加至所述文本信息中。

8.根据权利要求 1 至 7 任意一项所述的方法，还包括：

采用序列标注方法，对训练文本信息中的元素进行标注，得到每一个所述训练文本信息的标签；和

根据所述训练文本信息以及所述训练文本信息的标签，对所述目标模型进行训练。

9.根据权利要求 1 至 8 任意一项所述的方法，其中，所述关键词类别包括时间关键词、地点关键词、人物关键词、事件关键词中的至少一种。

10.一种识别关键词类别的装置，包括：

第一获取模块，用于获取待识别的文本信息；和

识别模块，用于将所述文本信息输入至目标模型，输出所述文本信息中的各个元素所属的关键词类别；

其中，所述目标模型包括用于提取所述文本信息的语义特征的第一处理层、用于提取所述文本信息中存在于业务词典中的目标关键词的第二处理层、以及用于基于所述语义特征和所述目标关键词确定所述文本信息中的各个元素所属的关键词类别的第三处理层，所述第一处理层与所述第二处理层并行设置；

所述业务词典包括在业务过程中已识别出的多个关键词。

11.一种电子设备，包括存储器、收发机和处理器：

所述存储器用于存储计算机程序；所述收发机用于在所述处理器的控制下收发数据；所述处理器用于读取所述存储器中的计算机程序并执行如权利要求 1 至 9 中任一项所述的识别关键词类别的方法。

12.一种处理器可读存储介质，所述处理器可读存储介质存储有计算机程序，所述计算机程序用于使处理器执行如权利要求 1 至 9 中任一项所述的识别关键词类别的方法。

13.一种计算机程序，包括：

指令，所述指令当由处理器执行时使所述处理器执行如权利要求 1 至 9 中任一项所述的识别关键词类别的方法。

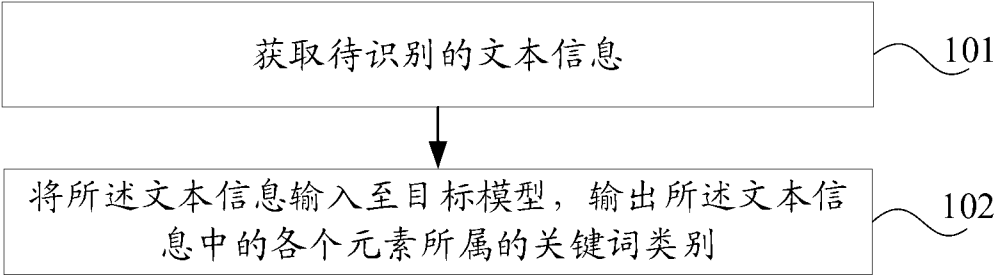


图 1

	外	卖		你	好	.	我	是	A	B	外	卖	.	请	下	来	取	您	的	快	递	.
外																						
卖																						
。																						
你																						
好																						
.																						
我																						
是																						
A																						
B																						
外																						
卖																						
.																						
请																						
下																						
来																						
取																						
您																						
的																						
快																						
递																						
。																						

图 2

	外	卖	。	你	好	，	我	是	A	B	外	卖	，	请	下	来	取	您	的	快	递	。
外																						
卖																						
。																						
你																						
好																						
，																						
我																						
是																						
A																						
B																						
外																						
卖																						
，																						
请																						
下																						
来																						
取																						
您																						
的																						
快																						
递																						
。																						

图 3

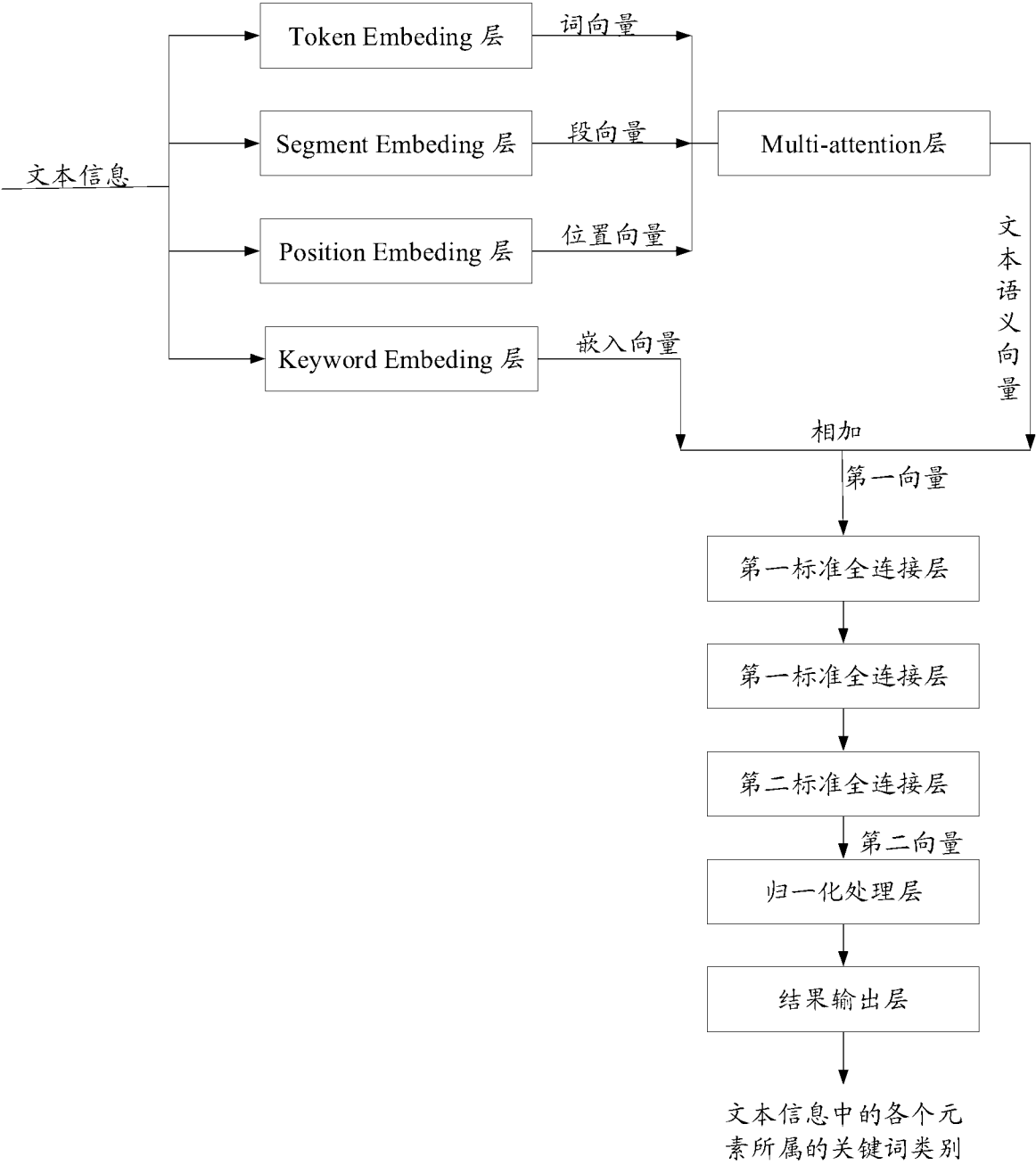


图 4



图 5



图 6

INTERNATIONAL SEARCH REPORT

International application No.

PCT/CN2023/136792

A. CLASSIFICATION OF SUBJECT MATTER

G06F 16/35(2019.01)i; G06F 40/289(2020.01)i; G06F 40/30(2020.01)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)

IPC: 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)

CNABS, CNTXT, CNKI, VEN, USTXT, EPTXT, WOTXT, IEEE: 关键词, 文本, 识别, 类别, 分类, 词典, 语义, 注意力, 掩码, 词, 段, 位置, keyword, text, recognition, classify, category, dictionary, semantic, attention, mask, token, segment, position

C. DOCUMENTS CONSIDERED TO BE RELEVANT

Category*	Citation of document, with indication, where appropriate, of the relevant passages	Relevant to claim No.
PX	CN 117076666 A (CHINA TELECOM CORP., LTD.) 17 November 2023 (2023-11-17) description, paragraphs 32-166	1-13
X	CN 112632292 A (SHENZHEN ONECONNECT SMART TECHNOLOGY CO., LTD.) 09 April 2021 (2021-04-09) description, paragraphs 73-113	1-13
A	CN 113886577 A (RUNLIAN SOFTWARE SYSTEM (SHENZHEN) CO., LTD.) 04 January 2022 (2022-01-04) entire document	1-13
A	CN 114781366 A (GUANGDONG OPPO MOBILE COMMUNICATIONS CO., LTD.) 22 July 2022 (2022-07-22) entire document	1-13
A	US 2023015606 A1 (TENCENT TECHNOLOGY (SHENZHEN) CO., LTD.) 19 January 2023 (2023-01-19) entire document	1-13



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

“D” document cited by the applicant in the international application

“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

01 March 2024

Date of mailing of the international search report

06 March 2024

Name and mailing address of the ISA/CN

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

Authorized officer

Telephone No.

INTERNATIONAL SEARCH REPORT
Information on patent family members

International application No.

PCT/CN2023/136792

Patent document cited in search report			Publication date (day/month/year)	Patent family member(s)			Publication date (day/month/year)
CN	117076666	A	17 November 2023	None			
CN	112632292	A	09 April 2021	None			
CN	113886577	A	04 January 2022	None			
CN	114781366	A	22 July 2022	None			
US	2023015606	A1	19 January 2023	WO	2022078102	A1	21 April 2022

A. 主题的分类

G06F 16/35(2019.01)i; G06F 40/289(2020.01)i; G06F 40/30(2020.01)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)

IPC: 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))

CNABS, CNTXT, CNKI, VEN, USTXT, EPTXT, WOTXT, IEEE: 关键词, 文本, 识别, 类别, 分类, 词典, 语义, 注意力, 掩码, 词, 段, 位置, keyword, text, recognition, classify, category, dictionary, semantic, attention, mask, token, segment, position

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 117076666 A (中国电信股份有限公司) 2023年11月17日 (2023 - 11 - 17) 说明书第32-166段	1-13
X	CN 112632292 A (深圳壹账通智能科技有限公司) 2021年4月9日 (2021 - 04 - 09) 说明书第73-113段	1-13
A	CN 113886577 A (润联软件系统(深圳)有限公司) 2022年1月4日 (2022 - 01 - 04) 全文	1-13
A	CN 114781366 A (OPPO广东移动通信有限公司) 2022年7月22日 (2022 - 07 - 22) 全文	1-13
A	US 2023015606 A1 (TENCENT TECHNOLOGY (SHENZHEN) CO., LTD.) 2023年1月19日 (2023 - 01 - 19) 全文	1-13

☐ 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

“D” 申请人在国际申请中引证的文件

“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

2024年3月1日

国际检索报告邮寄日期

2024年3月6日

ISA/CN的名称和邮寄地址

中国国家知识产权局
中国北京市海淀区蓟门桥西土城路6号 100088

授权官员

杜军

电话号码 (+86) 010-53961418

PCT/ISA/210 表(第2页) (2022年7月)

国际检索报告
关于同族专利的信息

国际申请号
PCT/CN2023/136792

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	117076666	A	2023年11月17日	无			
CN	112632292	A	2021年4月9日	无			
CN	113886577	A	2022年1月4日	无			
CN	114781366	A	2022年7月22日	无			
US	2023015606	A1	2023年1月19日	WO	2022078102	A1	2022年4月21日