

一种元搜索引擎的查询结果处理模型

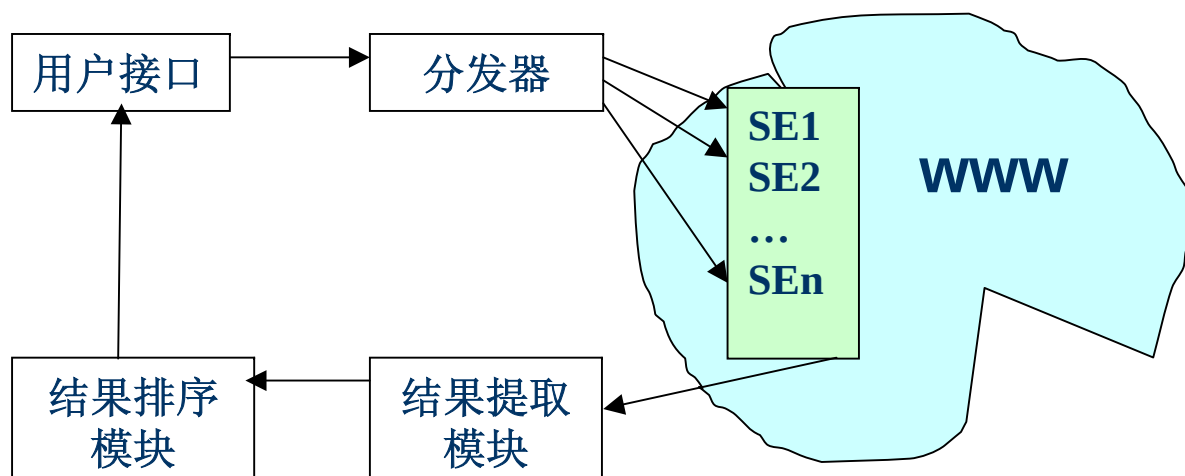
演讲：张强弓
北京工业大学



元搜索引擎

- 调用其它搜索引擎的搜索引擎
- 无需维护庞大的数据库
- 搜索的查全率高

元搜索引擎的体系结构

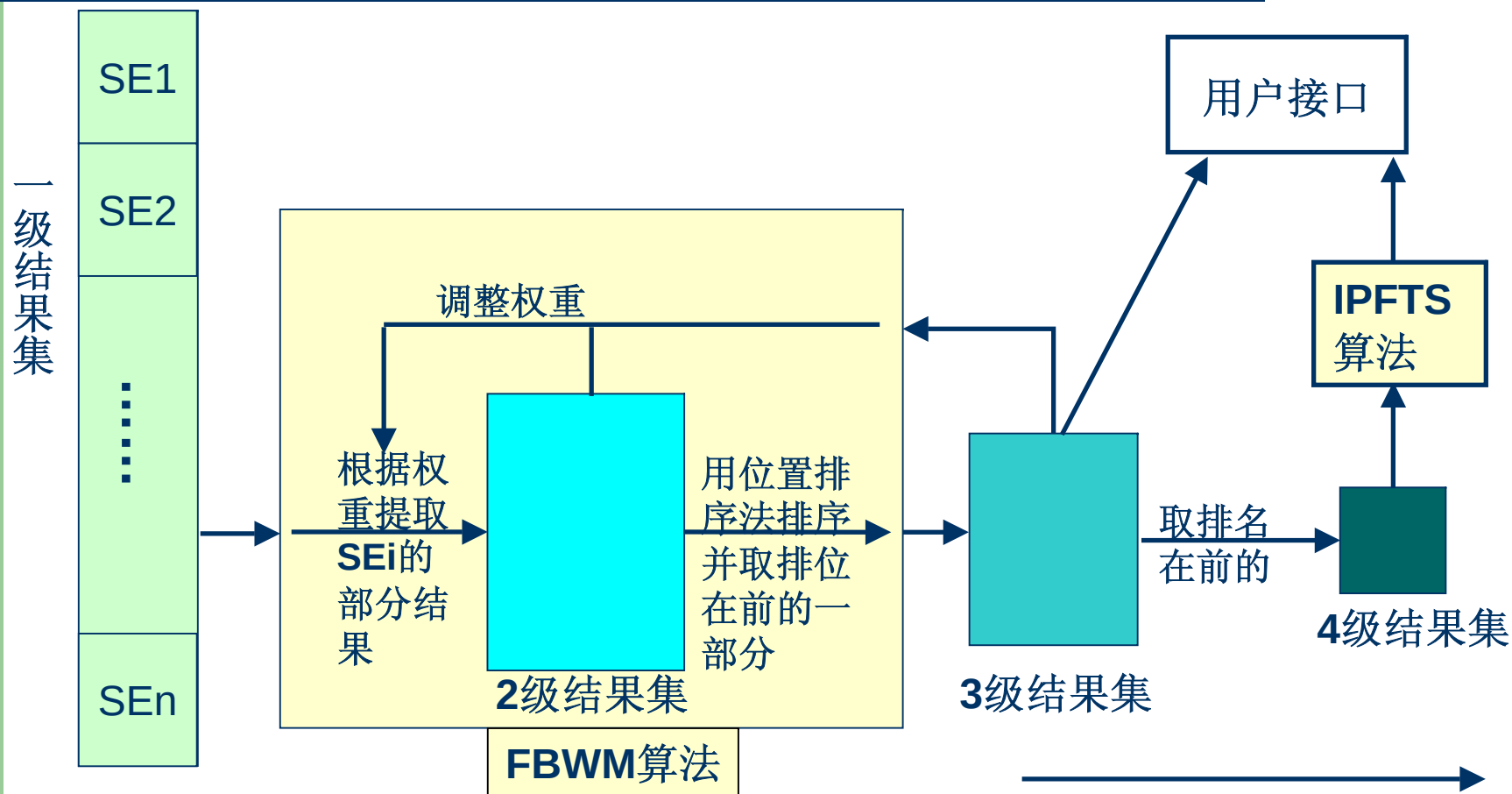


元搜索引擎的体系结构

本文中的查询结果处理模型

- 指的是查询结果的提取和查询结果的排序
- 关键部分：
 1. 四级结果集结构
 2. 根据反馈调整权重（**FBWM**）
 3. 改进的位置/全文排序（**IPFTS**）

模型的工作过程



1. 四级结果集结构
2. 根据反馈调整权重 (**FBWM**)
3. 改进的位置/全文排序 (**IPFTS**)

1.四级结果集结构的思想

- 元搜索引擎要处理的是非常庞大的结果集
- 我们需要精度更高的算法，但这会耗费更多的时间
- 于是，设置多级结果集，在庞大的较低级结果集中用低级快速的算法挑选出一部分好的结果构成较高级的结果集，在较高级的结果集中应用较高级的算法，形成更高级的结果集。如此往复。
- 目标是保证排名在前的少数结果的准确度。
- 依据是：用户最终可能浏览到的网页少于100个

2.FBWM算法(根据反馈调整权重)

- 结果提取部分
- 结果排序部分
- 调整权重部分

FBWM算法

- 符号的约定：
 - S_i 表示元搜索引擎调用的第 i 个独立搜索引擎，假设共调用 M 个独立搜索引擎；
 - R_i 表示第 i 个独立搜索引擎返回的结果的结果集；
 - W_i 为 S_i 的权重；
 - N_i 表示从 R_i 中根据 W_i 提取的结果的数量；
 - n_i 表示在3级结果集中，属于 R_i 中的结果的数量。

FBWM算法的描述

•结果提取部分

- ① 对每个独立搜索引擎 S_i 赋以权重 W_0 ，即 $W_i = W_0$
- ② 将每个 R_i 中前 N_i 个结果取出，合并，形成2级结果集，其中

$$N_i = c_1 * |R_i| * w_i / \sum_{i=1}^M w_i$$

$|R_i|$ 表示集合 R_i 的基数。 c_1 是常数，可以取0.1,0.01等。

FBWM算法的描述

- 结果排序部分

③对**2**级结果集用位置排序法进行排序，取出前**n**个结果形成**3**级结果集。其中：

$$n = c_2 * \sum_{i=1}^M N_i$$

n的大小由常数**C₂** 决定

FBWM算法的描述

- 调整权重部分

④ 对每个独立搜索引擎 S_i ，计算它对3级结果集的贡献比率 p_i ：

$$p_i = n_i / N_i$$

对 p_i 作规范化处理，得到贡献比率调节系数 P_i

$$P_i = p_i / \sum_{i=1}^M p_i$$

FBWM算法的描述

⑤ 重新调整每个 S_i 的权重 W_i :

$$W_i = C_3 * W_i + C_4 * P_i$$

其中 C_3 和 C_4 为常数。

对所有 W_i 重新计算完毕后，为了保证每次 P_i 对 W_i 的影响力是一样的，将每个 W_i 重新化成百分比形式：

$$w_i = w_i / \sum_{i=1}^n w_i$$

FBWM算法的描述

对每次查询，都重复步骤②到步骤⑤。

FBWM算法的特点

- 对搜索结果总数量的适应性
- 对独立搜索引擎性能变化的适应
- 对独立搜索引擎的搜索精度的差别的适应
 - 用贡献比率而不用贡献率
- 与 *训练集方法* 的不同之处

3.IPFTS方法（搜索结果的排序）

改进的位置/全文排序法



IPFTS方法的思想

- 位置/摘要排序¹法的改进
 - 摘要排序法简单速度快但准确度不高
 - 用全文信息代替摘要信息计算相关度
- 用词条匹配等级使相关度计算更加精确

¹ Refer to 徐宝文, 张卫峰 《搜索引擎与信息获取技术》

IPFTS方法

- 符号的约定

- r_i 表示第 i 个结果集中第 i 个结果;
- D_i 表示 r_i 对应的全部文本信息, 称作文档;
- q 表示用户输入的查询串;
- l_j 表示 q 中第 j 个词条。
- X 为查询串 q 中的词条数

词条匹配等级

- 假设一个查询串 q 有 m 个词条 $l_1, l_2, \dots, l_m$ ，如果在文档 D_1 中， l_1 出现了 m 次， $l_2, \dots, l_m$ 一次也没出现，而在长度与文档 D_1 相同的文档 D_2 中， $l_1, \dots, l_m$ 各出现一次。显然 D_2 是比 D_1 更好的结果。（例如：查询串 q = “信息 获取 技术”，则词条“信息”出现很多次而“获取”和“技术”没有出现的文档，与“信息”，“获取”，“技术”都出现的文档相比，自然是后者具有更高的相关度）
- 于是，我们考虑给词条匹配全面的文档奖励：更高的词条匹配等级。

词条匹配等级定义

- 定义：设查询串 q 有 X 个词条，文档 D 中包含有这 X 个词条的 N 个（ $N \leq X$ ），则文档 D 与查询串 q 的词条匹配等级为 N

IPFTS方法的描述

对第4结果集中的每个结果 r_i :

- ① 仿照摘要排序法，计算文档 D_i 与查询串 q 的普通相关度（是指没有词条匹配等级影响的相关度） $Rq(q, D_i)$

先计算每个词条与文档的相关度 $Rl(l_j, D_i)$

$$Rl(l_j, D_i) = \frac{Occurrence(l_j, D_i)}{\sum_{k=1} \ln(Length(D_i) / Location(l_j, k, D_i))}$$

再计算查询串与文档的相关度 $Rq(q, D_i)$

$$Rq(q, D_i) = \sum_j^X Rl(l_j, D_i)$$

IPFTS方法的描述

- ② 计算文档 D_i 与查询串 q 的词条匹配等级 $MG(q, D_i)$
- ③ 计算查询串 q 与文档 D_i 的最终相关度 $R(q, D_i)$:

$$R(q, D_i) = MG(q, D_i) * Rq(q, D_i)$$

- ④ 计算位置信息的得分 $P(r_i)$:

$$P(r_i) = 1 / i$$

IPFTS方法的描述

⑤ 综合位置信息和相关度信息得到最终排序分数 **Rank(r_i)**

$$Rank(r_i) = c_5 * R(q, D_i) / \sum_{i=1}^K R(q, D_i) + c_6 * P(r_i) / \sum_{i=1}^K P(r_i)$$

其中**C₅**，**C₆**是常数，**K**为第4结果集中结果的个数

对第4结果集中的结果按照**Rank(r_i)**值从大到小排序

IPFIS方法的特点

- 利用全文信息计算相关度能适应各个独立搜索引擎用户接口的差异，同时能保证排名在前的搜索结果的有效性（无死连接）
- 词条匹配等级的引入提高了相关度计算的准确性
- 位置排序法本身的特点

结束语

- 这个模型能减少元搜索引擎对独立搜索引擎的依赖，同时能提高元搜索引擎的查询速度。
- 将来如果能结合用户接口模块和分发器模块的新技术共同应用，将会使元搜索引擎更具威力。