

موسسه بابان

انتشارات بابان و انتشارات راهیان ارشد

پایگاه داده‌ها

فصل ششم: نرمال سازی (درسنامه)

(ویژه کنکور ارشد ۱۴۰۵)

ویژه‌ی داوطلبان کنکور کارشناسی ارشد مهندسی کامپیوتر و IT

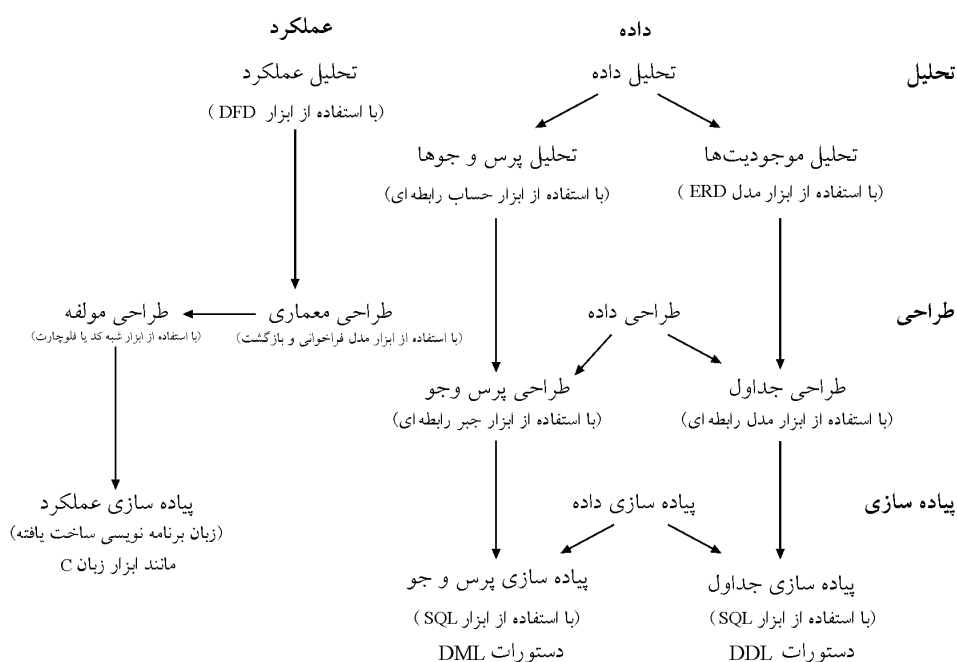
براساس کتب مرجع

آبراهام سیلبرشاتز، راما کریشنان، رامز المصری، سی جی دیت و توماس کانلی و کارولین بگ

ارسطو خلیلی فر

مقدمه

همانطور که در فصل وابستگی تابعی گفتیم، از نظر جانمایی مبحث **نرمال سازی** مربوط به همان فعالیت طراحی و مدل رابطه‌ای و پیش از شروع پیاده سازی است. چون جداول اگر غیرنرمال باشند، اول باید طراحی نرمال شوند، سپس پیاده سازی شوند تا از پیاده سازی غیرنرمال جلوگیری شود. شکل زیر گویای مطلب است:



جداول **غیرنرمال** به مشکل **افزودگی محتوایی (طبیعی یا داده‌ای)** مبتلا هستند و از آن **رنج**

می‌برند که خود این افزونگی ریشه در **تکرار اطلاعات** دارد و تکرار اطلاعات ریشه در وجود **وابستگی‌های بد** (افزونگی‌ساز و خرابکار) بین ستون‌های جدول دارد. راه **درمان** جداول **غیرنرمال** (**کثیف**) و تبدیل آن به جداول **نرمال** (**تمیز**)، فرآیند **نرمال‌سازی** است. به طور کلی سه نوع وابستگی بین ستون‌های جداول وجود دارد:

- **وابستگی تابعی (FD: Functional Dependency)**
- **وابستگی چند مقداری (MVD: Multi-Valued Dependency)**
- **وابستگی الحاقی (JD: Join Dependency)**

توجه: وابستگی چندمقداری (MVD) و وابستگی الحاقی (JD) نوع دیگری از وابستگی‌ها هستند که **تابعی نیستند**. شرح **خوب** و **بد** در ادامه همین بررسی می‌شود. به طور کلی دو نوع **وابستگی تابعی** بین ستون‌های جداول وجود دارد:

الف) وابستگی‌های تابعی خوب:

- ۱- **وابستگی تابعی ساده با ابرکلید (FD: Functional Dependency)**
 - ۲- **وابستگی تابعی کامل (تام) با کلید کاندید (FFD: Fully Functional Dependency)**
- ب) وابستگی‌های تابعی بد:

- ۱- **وابستگی تابعی بخشی (غیرکامل) (PFD: Partial Functional Dependency)**
- ۲- **وابستگی تابعی انتقالی (تراگذری) (TFD: Transitive Functional Dependency)**
- ۳- **وابستگی تابعی معکوس (RFD: Reverse Functional Dependency)**

در فصل سوم (مدل رابطه‌ای)، فرآیند تبدیل **مدل تحلیل** (نمودار ER) به **مدل طراحی** (جداول) را دیدیم و گفتیم که اگر این فرآیند **به‌درستی** انجام شود، نتیجه آن جداول **نرمال** (خوش طرح) است. اما اگر این فرآیند به درستی انجام نشود، نتیجه آن جداول **غیرنرمال** (**کثیف**) است که دچار **تکرار اطلاعات** و به تبع دچار **افزونگی داده‌ای** (**محتوایی یا طبیعی**) است و از آن **رنج** می‌برد که با عملیات **نرمال‌سازی** باید **درمان** شود. به مثال زیر دقت کنید:

مثال: جدول غیرنرمال studclg با مقادیر زیر را در نظر بگیرید:

وابستگی تابعی بد (از نوع وابستگی تابعی انتقالی)

S#	Sname	Clg#	Clgname	
8621	Ali	12	Computer	افزونگی طبیعی →
8442	Reza	12	Computer	
8731	Abbas	NULL	NULL	
9215	Hassan	13	Bargh	

اثرات وابستگی تابعی بد $Clg\# \rightarrow Clgname$ در جدول studclg :

۱- **افزونگی داده‌ای:** این اطلاعات که «کد دانشکده‌ی کامپیوتر برابر 12 است» در سطرهای اول و دوم تکرار شده است، پس **افزونگی داده‌ای (محتوایی یا طبیعی)** وجود دارد که باعث هدر رفتن حافظه شده است.

۲- **آنومالی:** اگر شخصی بخواهد کد دانشکده‌ی کامپیوتر را به 22 تغییر دهد و به اشتباه، فقط مقدار $Clg\#$ در سطر اول جدول studclg را عوض کند، از این به بعد دانشکده کامپیوتر دو کد خواهد داشت: 12 و 22، یعنی بالاخره کد دانشکده کامپیوتر 12 هست یا 22؟ این موضوع یعنی **آنومالی (بی‌نظمی، ناسازگاری، ناهنجاری)** در **آپدیت**. به این مدل ناسازگاری، ناسازگاری ارجاعی هم گفته می‌شود چون یک مقدار است که همه جا باید یکسان باشد ولی در اثر تغییر در جدول غیرنرمال همه جا یکسان نیست. پدیده ناسازگاری ارجاعی با نرمال سازی و تعریف کلید خارجی مرتفع می‌شود. **ناسازگاری داده‌ای** انواع مختلف دارد که در فصل مفاهیم پایه بررسی شد.

۳- **مقادیر NULL:** در سطر 3 از جدول studclg، مقادیر NULL وجود دارد. به طور کلی، به ازای هر سطر از جدول غیرنرمال studclg که $Clg\#$ آن برابر NULL باشد یک سطر در جدول studclg خواهیم داشت که مقادیر دو ستون آخر آن NULL خواهد بود. پس مشکل مقادیر NULL نیز وجود دارد.

توجه: برقراری وابستگی تابعی بد (از نوع وابستگی تابعی انتقالی) $Clg\# \rightarrow Clgname$ از جدول studclg به عنوان **محدودیت جامعیتی** به طور **صریح** توسط دستور **Assertion** و **Trigger** کنترل می‌شود.

انواع آنومالی در عملیات ذخیره سازی (درج، حذف و آپدیت)

- ۱- عدم امکان انجام یک عمل (که منطقاً باید قابل انجام باشد)
- ۲- بروز پیامد بد (عارضه جانبی) پس از انجام یک عمل
- ۳- بروز فزونکاری در سیستم در انجام یک عمل

آنومالی در درج: اگر رکورد (Computer , 22 , sajad , 5261) درج شود. از این به بعد دانشکده کامپیوتر دو کد خواهد داشت: 12 و 22 یعنی بالاخره کد دانشکده کامپیوتر 12 هست یا 22؟ این آنومالی از نوع «**بروز پیامد بد (عارضه جانبی) پس از انجام یک عمل**» است. برای مثالی دیگر اگر بخواهیم فقط رکورد (mathematic , 14) را برای واقعیت مشخصات دانشکده ریاضی درج کنیم، تا زمانی که $S\#$ برای آن مقدار ندارد، امکان درج رکورد (mathematic , 14) نیست چون $S\#$ کلید کاندید جدول غیرنرمال studclg است و محدودیت جامعیت موجودیت جلوی درج آنرا

می‌گیرد. این آنومالی از نوع «عدم امکان انجام یک عمل (که منطقاً باید قابل انجام باشد)» است. **آنومالی در حذف:** اگر قرار باشد فقط اطلاعات یک دانشکده (Clg#, Clgname) حذف شود در کنار آن، اطلاعات دانشجویی (S#, Sname) هم ناخواسته از دست می‌رود. این آنومالی از نوع «بروز پیامد بد (عارضه جانبی) پس از انجام یک عمل» است.

آنومالی در آپدیت: اگر کد دانشکده‌ی کامپیوتر به 22 تغییر کند و به اشتباه، فقط مقدار Clg# در سطر اول جدول studclg تغییر کند، از این به بعد دانشکده کامپیوتر دو کد خواهد داشت. بنابراین به طور دستی باید کد دانشکده‌ی کامپیوتر در همه سطرها به 22 تغییر کند. این آنومالی از نوع «بروز فزونکاری در سیستم در انجام یک عمل» است.

نتیجه: علت آنومالی‌های فوق، وجود وابستگی‌های تابعی بد در جدول studclg است، که با حذف وابستگی‌های بد، آنومالی در عملیات ذخیره‌سازی (درج، حذف و آپدیت) هم ایجاد نمی‌شود.

نتیجه: اگر فرآیند تبدیل مدل تحلیل (نمودار ER) به مدل طراحی (جداول) به درستی انجام نشود، نتیجه آن جداول غیرنرمال است که دچار مشکلات زیر است:

۱- افزونگی داده‌ای یا طبیعی یا محتوایی (Data Redundancy)

۲- آنومالی (مشکل در درج، حذف و آپدیت)

۳- مقادیر تهی (NULL Values)

نتیجه: برای مقدار NULL، فضا برای آن تخصیص می‌یابد اما برای هیچ و پوچ که باعث هدر رفتن حافظه می‌شود.

با تجزیه جدول studclg به دو جدول clg و stud، وابستگی تابعی بد Clg# → Clgname از جدول studclg حذف می‌شود و مشکل درمان می‌شود.

حال اگر جدول studclg را مطابق فرآیند نرمال‌سازی به دو جدول clg و stud تجزیه کنیم:

S#	Sname	Clg#	Clg#	Clgname
8621	Ali	12	12	Computer
8442	Reza	12	13	Bargh
8731	Abbas	NULL	جدول clg	
9215	Hassan	13		

جدول Stud

مشکلات فوق مرتفع و نتایج زیر حاصل خواهد شد:

۱ - کاهش افزونگی داده‌ای (طبیعی یا محتوایی) مانند تکرار بی دلیل سطر (12, Computer)

۲ - افزایش افزونگی تکنیکی (فنی یا ساختاری) به دلیل تعریف کلید خارجی

۳ - کاهش مقادیر NULL

۴ - عدم آنومالی (عدم مشکل در درج، حذف و آپدیت)

توجه: کاهش **بعضی** از انواع افزونگی در فرآیند نرمال سازی، یعنی افزونگی طبیعی کاهش می یابد، اما به واسطه عمل تجزیه و تعریف کلید خارجی (در قالب ستون های مشترک)، افزونگی تکنیکی (ساختاری) افزایش می یابد.

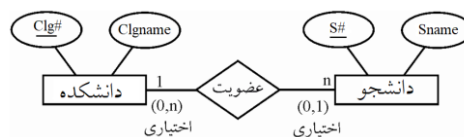
توجه: هر جا جدول دیدیم به آن نمی گوئیم پایگاه داده نرمال، بلکه هر جا جدول نرمال دیدیم به آن می گوئیم پایگاه داده نرمال.

توجه: دقت کنید نرمال سازی زمانی کاربرد دارد که خود مدل رابطه ای در سطح مدل طراحی به شکل بهینه انجام نشود، برای مثال رابطه یک به چند بین دو موجودیت در یک ERD یک مدل رابطه ای **بهینه** دارد. که اگر عمل نگاشت مدل تحلیل به مدل طراحی بهینه انجام شود جداول حاصل خود به خود نرمال هستن و نیاز به نرمال سازی ندارد.

نگاشت رابطه یک به چند بین دو موجودیت به مدل رابطه ای

مستقل از اختیاری یا اجباری بودن موجودیت ها، هر موجودیت به یک جدول تبدیل می گردد. و کلید کاندید جدول یک در جدول چند به عنوان کلید خارجی تعریف می گردد.

مدل تحلیل:



مدل طراحی درست:

<u>S#</u>	Sname	Clg#
8621	Ali	12
8442	Reza	12
8731	Abbas	NULL
9215	Hassan	13

جدول Stud

<u>Clg#</u>	Clgname
12	Computer
13	Bargh

جدول clg

توجه: اگر دانشجویی عضو دانشکده‌ای نباشد، ستون Clg# برای آن مقدار NULL می‌گیرد، مانند شماره دانشجویی 8731 در جدول Stud. همچنین Clg# در جدول Stud **کلید خارجی** است.

توجه: همانطور که مشاهده می‌کنید اگر اصول تحلیل (مدل ER) و طراحی (مدل رابطه‌ای) رعایت شود. جداول حاصل، خود به خود نرمال هستند و نیاز به نرمال‌سازی ندارد.

بررسی سربار سیستم (Overhead) در عمل بازیابی قبل و بعد از نرمال‌سازی

جدول غیرنرمال studclg به صورت زیر است:

S#	Sname	Clg#	Clgname
8621	Ali	12	Computer
8442	Reza	12	Computer
8731	Abbas	NULL	NULL
9215	Hassan	13	Bargh

جدول studclg

جداول نرمال stud و clg نیز به صورت زیر است:

S#	Sname	Clg#	Clg#	Clgname
8621	Ali	12	12	Computer
8442	Reza	12	13	Bargh
8731	Abbas	NULL		
9215	Hassan	13		

جدول clg

جدول Stud

پرس‌وجو ۱: (بعد از نرمال‌سازی)

select S#

from stud

ستون S# از جدول نرمال stud استخراج می‌شود و نیازی به الحاق دو جدول stud و clg ندارد.

پرس‌وجو ۲: (قبل از نرمال‌سازی)

select S#

from studclg

ستون S# از جدول غیرنرمال studclg استخراج می‌شود، اما با زمان واکنشی **بیشتر** نسبت به جدول

نرمال stud، چون طول رکوردها در جدول غیرنرمال studclg بیشتر از جدول نرمال stud است. **نتیجه:** اگر پرس وجو فقط وابسته به یکی از جداول نرمال باشد مثل ستون S# از جدول stud، در این شرایط واکنشی S# از جدول stud زمان کمتری نسبت به واکنشی S# از جدول studclg دارد، چون در نرخ انتقال برابر، طول رکورد جدول stud از جدول studclg کمتر است.

$$\text{طول رکورد} \left(\overline{T_f} \right) = \frac{L}{R} \quad \text{نرخ انتقال}$$

پرس وجو ۳: (بعد از نرمال سازی)

```
select *
from stud, clg
where stud.clg# = clg.clg# AND clgname = 'computer'
```

ستونهای مشخصات دانشجویان دانشکده کامپیوتر از دو جدول نرمال Stud و clg استخراج می شود و نیاز به الحاق دو جدول stud و clg دارد، بنابراین سربرار حاصل از الحاق دو جدول را هم دارد.

پرس وجو ۴: (قبل از نرمال سازی)

```
select *
from studclg
where clgname = 'computer'
```

ستونهای مشخصات دانشجویان دانشکده کامپیوتر از جدول غیرنرمال studclg استخراج می شود، اما با زمان واکنشی **کمتر** نسبت به الحاق جداول نرمال stud, clg، هرچندکه طول رکوردها در جدول studclg **برابر** الحاق جداول stud, clg است. اما در حالت جداول نرمال، برای داشتن تمام ستونها نیاز به یک عمل اضافه و دارای سربرار به نام **الحاق** وجود دارد.

نتیجه: نرمال سازی باعث تجزیه جداول می شود، اگر پرس وجو روی جداول نرمال نیاز به الحاق نداشته باشد، سربرار حاصل از الحاق ایجاد نمی شود، که منجر به **کاهش** زمان پاسخگویی نسبت به جدول غیرنرمال می شود. اما اگر پرس وجو روی جداول نرمال نیاز به الحاق داشته باشد، سربرار حاصل از الحاق ایجاد می شود، که منجر به **افزایش** زمان پاسخگویی نسبت به جدول غیرنرمال می شود.

فرآیند نرمال سازی

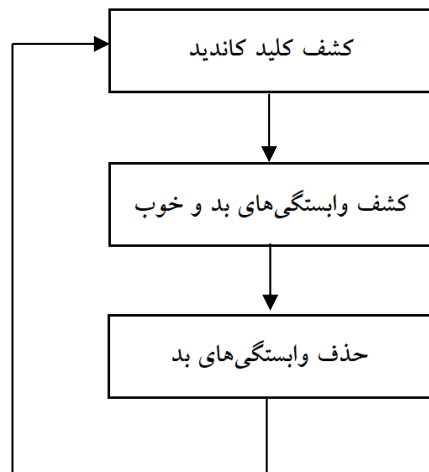
فرآیند نرمال سازی یعنی **تبدیل** یک جدول پایه غیرنرمال به چند جدول نرمال از طریق **تجزیه عمودی یا پارسازی عمودی (Vertical DE composition)**. در این فرآیند، انواع **وابستگی های بد** به ترتیب، سطح به سطح و به طور سلسله مراتبی (و نه یکباره) **حذف** می شوند که به تبع منجر به کاهش تکرار اطلاعات و کاهش افزونگی دادهای به طور سلسله مراتبی می شود. فرآیند

نرمال سازی از سه قدم زیر تشکیل شده است:

۱- کشف کلید کاندید

۲- کشف وابستگی های بد و خوب

۳- حذف وابستگی های بد

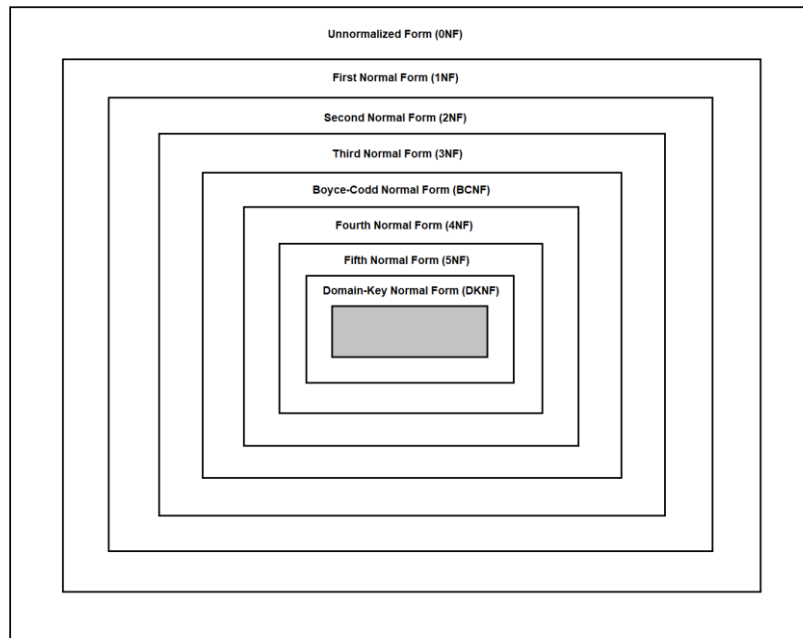


توجه: در دیاگرام فوق، چرخه به این دلیل است که فرآیند نرمال سازی، عملی به یکباره نیست، بلکه وابستگی های بد، سطح به سطح و به طور سلسله مراتبی در یک **رویه کاهش** از جداول حذف می شوند، در نتیجه افزونگی داده ای نیز سطح به سطح کاهش می یابد و نه یکباره.

توجه: فرآیند نرمال سازی یک **رویه بازسازنده (بازگشت پذیر)** است، یعنی از جداول نرمال می توان جدول پایه (اولیه) غیرنرمال را به طور **بی کاست** (بدون کم شدن سطر و ستون) و **بی حشو** (بدون اضافه شدن سطر و ستون) مجدد به دست آورد. اگر از همان راه که آمدیم، از همان راه هم برگردیم، مجدد به همان نقطه مبدا می رسیم.

سطوح نرمال سازی

سطوح نرمال سازی دارای خواص سلسله مراتبی است، بدین معنی که چنانچه جدولی در سطح k ام نرمال قرار داشته باشد، حتما در سطح $k-1$ نرمال نیز قرار دارد. بدیهی است که با افزایش سطح نرمال، احراز شرایط آن سطح نیز مشکل تر خواهد بود. همچنین هر سطح از سطح قبلی خود نرمال تر است. شکل زیر گویای مطلب است:



توجه: شکل فوق نشان می دهد، هر سطح داخلی زیر مجموعه سطح خارجی است:

$$5NF \subset 4NF \subset BCNF \subset 3NF \subset 2NF \subset 1NF \subset 0NF$$

مثال: 2NF ها از مجموعه 1NF ها هستند که به 2NF هم رسیدن.

توجه: راه در جلوی نرمال سازی باز است، یعنی همیشه می توان، جدولی که در یک سطح از نرمال قرار دارد، در سطح بالاتر نرمال آنرا قرار داد.

توجه: فرآیند نرمال سازی همانند فرآیند ادامه تحصیل (مقدماتی (0NF)، دیپلم (1NF)، کاردانی (2NF)، کارشناسی (3NF)، کارشناسی ارشد (BCNF)، دکتری (4NF) و پست دکتری (5NF)) است و همچنین دارای خواص سلسله مراتبی است، بدین معنی که چنانچه فردی در مقطع دکتری باشد، حتما کارشناسی ارشد و به قبل هم دارد، اما عکس آن صادق نیست، یعنی اگر شخصی کارشناسی ارشد داشته باشد، دلیل نمی شود که حتما دکتری هم دارد. همچنین، شخصی که کارشناسی دارد، همیشه راه برای ادامه تحصیل آن تا سطوح بالاتر وجود دارد.

توجه: به سطوح نرمال، **صورت نرمال** و **دیس بهنجار** هم گفته می شود.

شروط قرار داشتن یک جدول در یک سطح از نرمال

شرط لازم: در نرمال سطح قبلی قرار داشته باشد.

شرط کافی: قواعد تعریف شده در آن سطح از نرمال را پاس کرده باشد.

مثال: یک شخص فارغ التحصیل مقطع دکتری است:

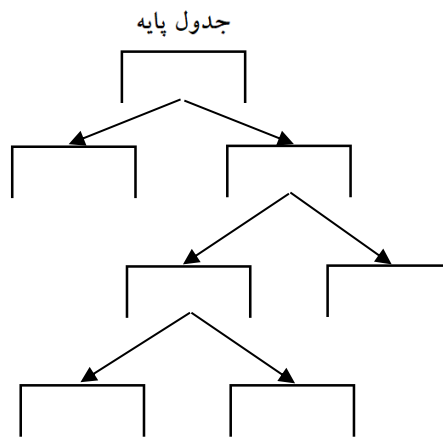
شرط لازم: کارشناسی ارشد داشته باشد.

شرط کافی: شرایط مقطع دکتری را به طور کامل پاس کرده باشد.

توجه: فرآیند نرمال سازی، مبتنی بر **تجزیه** جداول است، و حاصل آن یک **درخت تجزیه** است که

برگ های آن جداول نرمال (که همگی در سطح k ام نرمال قرار دارند) و **ریشه آن جدول پایه**

غیرنرمال است.



توجه: از نظر تئوری، تعداد سطوح نرمال، می تواند تا هر سطحی ادامه یابد، اما از نظر عملی در

اغلب موارد رسیدن به سطح سوم نرمال (3NF) برای یک جدول لازم و کافی است.

مشکلات پایگاه داده غیر نرمال

هر جدول ممکن است مشکلات مختلفی داشته باشد. در جدول زیر مشکلات پایگاه داده آمده است:

شماره	مشکل	راه حل
1	وجود صفت چندمقداری و مرکب	نرمال فرم اول (1NF)
2	وجود وابستگی تابعی بخشی	نرمال فرم دوم (2NF) برای حذف وابستگی بخشی
3	وجود وابستگی تابعی انتقالی	نرمال فرم سوم (3NF) برای حذف وابستگی انتقالی
4	وجود وابستگی تابعی معکوس	نرمال فرم (BCNF) برای حذف وابستگی معکوس
5	وجود وابستگی چندمقداری	نرمال فرم چهارم (4NF) برای حذف وابستگی چندمقداری
6	وجود وابستگی الحاقی	نرمال فرم پنجم (5NF) برای حذف وابستگی الحاقی

نرمال فرم اول (First Normal Form-1NF)

فرآیند نرمال سازی یک روند سطح به سطح است. سطح **نخست** نرمال سازی، تبدیل جداول به **نرمال فرم اول** است.

شروط قرار داشتن یک جدول در نرمال فرم اول به صورت زیر است:

شرط لازم: در نرمال سطح قبلی قرار داشته باشد، که قرار دارد چون نقطه شروع 0NF است.

شرط کافی: قواعد تعریف شده در آن سطح از نرمال را پاس کرده باشد، که به صورت زیر است:

۱ - دارای حداقل یک کلید کاندید باشد.

۲ - همه ستون های آن غیرقابل تجزیه و ساده باشند (جدول باید فاقد ستون مرکب باشد)

۳ - همه ستون های آن تکمقداری باشند (جدول باید فاقد ستون چندمقداری باشد)

توجه: شرط کافی، همان قانون **جامعیت درون رابطه ای** (Intra-Relational Integrity Rule) است، که مطابق این قانون، هر رابطه (جدول) باید به تنهایی درست باشد. یعنی صفات رابطه

اتومیک باشد (چند مقداری و مرکب نباشد) و همچنین هر رابطه باید حتماً حداقل دارای یک کلید کاندید باشد.

توجه: در یک رابطه، تمام مولفه‌ها در تاپل‌ها تجزیه‌ناپذیرند. به عبارت دیگر در مدل رابطه‌ای نمی‌توان فیلد مرکبی که از فیلدهای ساده تشکیل شده‌است، تعریف نمود. یعنی فیلدهای موجود در جداول مدل رابطه‌ای **نباید مرکب** مثل تاریخ تولد و آدرس باشد و **نباید چندمقداری** مثل مدرک تحصیلی و شماره تلفن باشد. به عبارت دیگر هیچ رابطه‌ای نمی‌تواند شامل تاپل‌های تودرتو، تاپل تو تاپل یا جدول تو جدول باشد، به بیان دیگر مولفه‌های هریک از تاپل‌های یک رابطه نباید مرکب یا چندمقداری باشد، به عبارت ساده‌تر هر مولفه یک تاپل، نمی‌تواند خودش یک تاپل باشد.

مثال: مجموعه R را در نظر بگیرید:

$$R = \{((a, x), 1), (a, x, 2), (b, x, 1), (b, x, 2)\}$$

مجموعه R یک رابطه نیست، زیرا مولفه نخست در تاپل اول یعنی (a, x) از $((a, x), 1)$ خودش یک تاپل است که مطابق قوانین مدل رابطه‌ای، هر مولفه در یک تاپل، نمی‌تواند خودش یک تاپل باشد.

مثال ۱ (نرمال‌سازی صفت مرکب تاریخ تولید):

طراحی نادرست (غیرنرمال)

$$R = \{(S1, Sn1, (1373, 7, 5)), (S2, Sn2, (1381, 6, 4))\}$$

SID	Sname	Date of birth
S1	Sn1	1373,7,5
S2	Sn2	1381,6,4

جدول غیرنرمال R

طراحی درست (نرمال) (بدون تجزیه جدول پایه)

$$R = \{(S1, Sn1, 1373, 7, 5), (S2, Sn2, 1381, 6, 4)\}$$

SID	Sname	Year	Month	Day
S1	Sn1	1373	7	5
S2	Sn2	1381	6	4

جدول نرمال R

نتیجه: جدول نرمال R در **نرمال فرم اول** قرار دارد، چون توسط نرمال‌سازی سه ستون Year، Month و Day برای آن تعریف شده‌است و از فرم تاپل تو تاپل خارج شده است. در روند فوق، نرمال‌سازی **بدون تجزیه** جدول پایه، و صرفاً فقط توسط **درج ستون** انجام شده است. کلید

کاندید جدول نرمال R ستون SID است.

مثال ۲ (نرمال سازی صفت چندمقداری شماره تلفن):

طراحی نادرست (غیرنرمال)

$R = \{(S1, Sn1, (22445566, 88774411)), (S2, Sn2, (44332211, NULL))\}$

وابستگی بد (از نوع وابستگی چندمقداری)

SID		Tel
Sname		
S1	Sn1	22445566,88774411
S2	Sn2	44332211,NULL

جدول غیرنرمال R

طراحی درست (نرمال) (بدون تجزیه جدول پایه)

$R = \{(S1, Sn1, 22445566, 88774411), (S2, Sn2, 44332211, NULL)\}$

SID	Sname	Tel1	Tel2
S1	Sn1	22445566	88774411
S2	Sn2	44332211	NULL

جدول نرمال R

نتیجه: جدول نرمال R در **نرمال فرم اول** قرار دارد، چون توسط نرمال سازی دو ستون Tel1 و Tel2 برای آن تعریف شده است و از فرم تاپل تو تاپل خارج شده است. در روند فوق، نرمال سازی **بدون تجزیه** جدول پایه، و صرفاً فقط توسط **درج ستون** انجام شده است. کلید کاندید جدول نرمال R ستون SID است.

طراحی درست (نرمال) (با تجزیه جدول پایه)

با **تجزیه جدول**، وابستگی بد (از نوع وابستگی چندمقداری) $SID \rightarrow \rightarrow Tel$ از جدول غیرنرمال

R **حذف** می شود و مشکل **درمان** می شود. به صورت زیر:

وابستگی کامل

SID	Sname
S1	Sn1
S2	Sn2

R1

SID	Tel
S1	22445566
S1	88774411
S2	44332211

R2

نتیجه: هر دو جدول R1 و R2 در **نرمال فرم اول** قرار دارند، چون جدول غیرنرمال R توسط نرمال سازی به دو جدول R1 و R2 **تجزیه** شد و از فرم تاپل تو تاپل خارج شده است. کلید کاندید

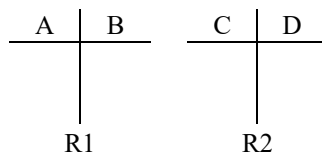
جدول R1 ستون SID و کلید کاندید جدول R2 تمام کلید (ستونهای SID, Tel باهم) است.

شروط تجزیه مطلوب در نرمال سازی

سوال: آیا روال تجزیه دلخواهی است؟ هر جوری دلم بخواهد است؟

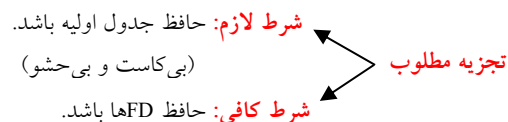
پاسخ: خیر، روال تجزیه دلخواهی نیست، بلکه برای رسیدن به تجزیه مطلوب، باید **شروط تجزیه مطلوب** محقق شود، وگرنه **تجزیه نامطلوب** می شود، که کاربردی هم ندارد.

مثال: رابطه $R(A, B, C, D)$ را با چهار ستون در نظر بگیرید، تجزیه زیر برای آن **تجزیه نامطلوب** است، چون بین دو جدول R1 و R2 **ستون مشترک** وجود ندارد و به تبع ارتباط بین دو جدول برای بازسازی **جدول اولیه** قطع می شود.



نتیجه: تجزیه مطلوب باید از دواج پذیر (پیوندپذیر) باشد، یعنی بین دو جدول، **ستون مشترک** وجود داشته باشد.

در فرآیند نرمال سازی شروط **تجزیه مطلوب** به صورت زیر است:



توجه: همانطور که قبل تر گفتیم، **نرمال سازی** یک **رویه بازسازنده (بازگشت پذیر)** است، یعنی از جداول تجزیه شده نرمال، باید بتوان جدول اولیه (پایه) غیرنرمال را به دست آورد.

واضح است که $R = R1 \text{ JOIN } R2$ برقرار است و جدول اولیه **دقیقا** به دست می آید، یعنی الحاق (پیوند) بازسازنده (بازگشت پذیر) است. در ادامه به شرح شرط لازم و کافی **تجزیه مطلوب** می پردازیم:

۱- شرط لازم: (تجزیه، حافظ تمام ستون ها (ساختار) و سطرها (محتوا) جدول اولیه باشد).
Non-Loss (LossLess) DE composition باشد، یعنی **تجزیه بی کاست** (بی گمشدگی) باشد.
Non additive DE composition باشد، یعنی **تجزیه بی حشو** (بی اضافه شدگی) باشد.
توجه: تجزیه مطلوب، باید به طور هم زمان هر دو شرط **بی کاست** و **بی حشو** را داشته باشد.
بی کاست (بدون کم شدن سطر و ستون نسبت به جدول اولیه): یعنی سطر و ستونی نسبت به جدول اولیه **کم نشود**، این مورد زمانی محقق می شود که **صفت مشترک** میان دو جدول وجود داشته باشد، اما کافی نیست چون فقط امکان بازسازی رکوردهای اولیه را دارد و ممکن است رکوردهای اولیه با رکوردهای اضافه تری بازسازی شود. بنابراین فقط بی کاست بودن از دقت کافی برخوردار نیست.
بی حشو (بدون اضافه شدن سطر و ستون جدید نسبت به جدول اولیه): یعنی سطر و ستونی نسبت به جدول اولیه **اضافه نشود**، این مورد زمانی محقق می شود که **صفت مشترک** در دو جدول، بر اساس بستار وابستگی های تابعی، **حداقل** در یکی از دو جدول **کلید کاندید** باشد.
تجزیه بی حشو به این معناست که از پیوند دو جدول تجزیه شده تاپل های حشو (تاپل های اضافی یا تاپل های ساختگی که در رابطه اولیه وجود نداشته است) پدید نیاید. در واقع اگر از پیوند دو جدول تجزیه شده، تاپل های حشو پدید آید، دیگر اطمینان نداریم که چه تاپل هایی اصلی و چه تاپل هایی اضافی هستند، بنابراین اطلاعات درست در بازسازی از دست می رود. نتیجه اینکه بی حشو بودن از بی کاست بودن مهم تر است، ضمن اینکه بی حشو بودن، بی کاست بودن هم پوشش می دهد. **پیوند بی حشو بودن**، از نوع **CK-CK** یا **CK-FK** است، که در هیچ یک از این دو نوع پیوند، تاپل حشو (اضافه) بروز نمی کند، یعنی بازسازنده است و تجزیه بی حشو می شود. در واقع تجزیه باید به گونه ای انجام شود که امکان **بازسازی ساختاری (ستونی)** و **بازسازی محتوایی (سطری) جدول اولیه** براساس جداول تجزیه شده امکان پذیر باشد، به این صورت که با الحاق جداول تجزیه شده باهم، جدول اولیه **دقیقا** به دست آید، به طوری که سطر و ستونی از جدول اولیه از دست نرود و همچنین سطر و ستونی به جدول اولیه اضافه نشود، **نه چیزی کم بشه نه چیزی اضافه بشه**. نتیجه اینکه در تجزیه مطلوب، هیچ اطلاعات سطری و ستونی نسبت به جدول اولیه کم و زیاد نمی شود، چون با پیوند مجدد جداول تجزیه شده، جدول اولیه به طور کامل در دسترس

است.

نتیجه: برای تحقق اینکه نه چیزی کم شود و نه چیزی اضافه شود، باید **صفت مشترک** در دو جدول، بر اساس بستار وابستگی های تابعی **حداقل** در یکی از دو جدول تجزیه شده **کلید کاندید** باشد. به عبارت دیگر اگر تجزیه به گونه ای باشد که **تعریف کلید خارجی** برقرار باشد، یعنی دو جدول تجزیه شده R_1 و R_2 **الحاق پذیر** باشند و حتما **صفت مشترک** در دو جدول، **حداقل** در یکی **کلید کاندید** و در دیگری **کلید خارجی** باشد، آنگاه شرط لازم **بی کاست** و **بی حشو** باهم محقق شده است.

نکته: اگر جدول اولیه R به دو جدول R_1 و R_2 تجزیه **بی کاست** و **بی حشو** شود، در این حالت $R = R_1 \text{ JOIN } R_2$ برقرار است و جدول اولیه به دست می آید، یعنی الحاق (پیوند) بازسازنده (بازگشت پذیر) است.

چک ستونی تجزیه مطلوب

اینکه بعد از عمل تجزیه، **ستونی کم نشود** و **ستونی اضافه نشود** توسط عملگر اجتماع $\cup$ و به صورت زیر چک می شود:

$$R_1 \cup R_2 = R$$

مثال: رابطه $R(A, B, C)$ با مجموعه وابستگی های تابعی (FD) زیر مفروض است:

$$F = \{ A \rightarrow C \}$$

پس از تجزیه به دو رابطه $R_1(A, B)$ و $R_2(A, C)$ ، ستون **A** ستون مشترک بین دو جدول است و $R_1 \cup R_2 = R$ **چک** می کند **ستونی کم نشود** و **ستونی اضافه نشود**، به صورت زیر:

$$R_1 \cup R_2 = \{A, B\} \cup \{A, C\} = \{A, B, C\}$$

چک سطری تجزیه مطلوب

اینکه بعد از عمل تجزیه، **سطری کم نشود** و **سطری اضافه نشود** باید چک شود که **صفت مشترک** یعنی $R_1 \cap R_2$ ، **حداقل** در یکی از دو جدول تجزیه شده R_1 و R_2 ، **کلید کاندید** است یا خیر؟ این کار توسط عملگر اشتراک $\cap$ و به صورت زیر انجام می شود:

$$R_1(H) \cap R_2(H) = \text{C.K. of } R_1(H) \quad \text{OR} \quad R_1(H) \cap R_2(H) = \text{C.K. of } R_2(H)$$

$R_i(H)$: یعنی صفات رابطه R_i

مثال: رابطه $R(A, B, C)$ با مجموعه وابستگی های تابعی (FD) زیر مفروض است:

$$F = \{ A \rightarrow C \}$$

پس از تجزیه به دو رابطه $R1(A, B)$ و $R2(A, C)$ ، **ستون مشترک A** بین دو جدول در کدام، **کلید کاندید** است؟

چک جدول R1

$$1: R1(H) \cap R2(H) \rightarrow R1(H)$$

$$\{A, B\} \cap \{A, C\} \rightarrow \{A, B\}$$

$$\{A\} \rightarrow \{A, B\}$$

$$2: R1(H) \cap R2(H) \rightarrow R1(H) - R2(H)$$

$$\{A, B\} \cap \{A, C\} \rightarrow \{A, B\} - \{A, C\}$$

$$\{A\} \rightarrow \{B\}$$

$$\{A\}^+ = \{A, \text{Stop}\}$$

توضیح: مطابق وابستگی‌های تابعی، **صفت مشترک A** کلید کاندید جدول **R1** نیست، چون همه ستون‌های جدول R1 را تولید نمی‌کند، صفات AB کلید کاندید جدول R1 است.

چک جدول R2

$$1: R1(H) \cap R2(H) \rightarrow R2(H)$$

$$\{A, B\} \cap \{A, C\} \rightarrow \{A, C\}$$

$$\{A\} \rightarrow \{A, C\}$$

$$2: R1(H) \cap R2(H) \rightarrow R2(H) - R1(H)$$

$$\{A, B\} \cap \{A, C\} \rightarrow \{A, C\} - \{A, B\}$$

$$\{A\} \rightarrow \{C\}$$

$$\{A\}^+ = \{A\}$$

$$A \rightarrow C$$

$$\{A\}^+ = \{A, C\}$$

توضیح: مطابق وابستگی‌های تابعی، **صفت مشترک A** کلید کاندید جدول **R2** است، چون همه ستون‌های جدول R2 را تولید می‌کند.

توضیح: عبارت 1 و 2 در **متون کتب مرجع مختلف**، یک معنی و نتیجه یکسان دارند، در عبارت 1 **وابستگی های بدیهی** هم تولید می شود، درحالی که در عبارت 2 به واسطه تفاضل، وابستگی های بدیهی حذف می شود. $R1(H) \cap R2(H)$ یعنی **صفت مشترک**، عبارت $R1(H)$ یعنی همه ستون های $R1$ ، عبارت $R2(H)$ یعنی همه ستون های $R2$ ، عبارت $R1(H) - R2(H)$ یعنی همه ستون های $R1$ منهای صفات مشترک و عبارت $R2(H) - R1(H)$ یعنی همه ستون های $R2$ منهای صفات مشترک، هر دو عبارت 1 و 2 بررسی می کند که **صفت مشترک** در جدول مورد بررسی **کلید کاندید** است یا خیر؟

مثال: رابطه $R(A, B, C, D, E)$ با مجموعه وابستگی های تابعی (FD) زیر مفروض است:

$$F = \{ (C, D) \rightarrow E \}$$

پس از تجزیه به دو رابطه $R1(A, B, C, D)$ و $R2(A, C, D, E)$ ، صفات مشترک **ACD** بین دو جدول در کدام، **کلید کاندید** است؟

چک جدول R1

$$1: R1(H) \cap R2(H) \rightarrow R1(H)$$

$$\{A, B, C, D\} \cap \{A, C, D, E\} \rightarrow \{A, B, C, D\}$$

$$\{A, C, D\} \rightarrow \{A, B, C, D\}$$

$$2: R1(H) \cap R2(H) \rightarrow R1(H) - R2(H)$$

$$\{A, B, C, D\} \cap \{A, C, D, E\} \rightarrow \{A, B, C, D\} - \{A, C, D, E\}$$

$$\{A, C, D\} \rightarrow \{B\}$$

$$\{A, C, D\}^+ = \{A, C, D, \text{Stop}\}$$

توضیح: مطابق وابستگی های تابعی، **صفات مشترک ACD** کلید کاندید جدول **R1** نیست، چون همه ستون های جدول $R1$ را تولید نمی کند.

چک جدول R2

$$1: R1(H) \cap R2(H) \rightarrow R2(H)$$

$$\{A, B, C, D\} \cap \{A, C, D, E\} \rightarrow \{A, C, D, E\}$$

$$\{A, C, D\} \rightarrow \{A, C, D, E\}$$

$$2: R1(H) \cap R2(H) \rightarrow R2(H) - R1(H)$$

$$\{A, B, C, D\} \cap \{A, C, D, E\} \rightarrow \{A, C, D, E\} - \{A, B, C, D\}$$

$$\{A, C, D\} \rightarrow \{E\}$$

$$\{A, C, D\}^+ = \{A, C, D\}$$

$$CD \rightarrow E$$

$$\{A, C, D\}^+ = \{A, C, D, E\}$$

توضیح: مطابق وابستگی‌های تابعی، **صفات مشترک ACD** کلید کاندید جدول **R2** است، چون همه ستون‌های جدول **R2** را تولید می‌کند.

نتیجه: در جدول $R2(A, C, D, E)$ صفات **ACD** به طور **بدیهی**، خودش خودش را می‌دهد و **E** را هم می‌دهد، بنابراین صفات **ACD** کلید کاندید جدول **R2** است.

نکته بسیار مهم: اگر بین دو جدول تجزیه شده فقط و فقط **صفات مشترک** باشد و صفات مشترک **حداقل** در یکی از دو جدول **کلید کاندید** نباشد، این حالت فقط تضمین می‌کند که نسبت به جدول اولیه **سطری کم نشود**، اما تضمین نمی‌کند که **سطری اضافه نشود**، بنابراین برای اینکه به طور هم‌زمان **تضمین** شود که **سطری کم نشود** و **سطری اضافه نشود** باید **صفات مشترک** در **حداقل** یکی از دو جدول **کلید کاندید** باشد.

مثال: برای نمونه یک تجزیه نامطلوب با ستون مشترک **C** به صورت زیر است:

A	C	⋈	B	C	=	A	B	C
1	201		101	201		1	101	201
3	201		102	201		1	102	201
4	204		103	204		3	101	201
						3	102	201
						4	103	204
R1			R1			R		

واضح است که $R = R1 \text{ JOIN } R2$ برقرار نیست (سطر سوم **اضافه** است) و جدول اولیه **دقیقا** به دست نمی آید، یعنی الحاق (پیوند) بازسازنده (بازگشت پذیر) نیست. چون **ستون مشترک C** در **حداقل** یکی از دو جدول **کلید کاندید** نیست.

۲- شرط کافی: (تجزیه، حافظ تمام FDهای جدول اولیه باشد).

FDs preserving DE composition باشد، یعنی **تجزیه حافظ وابستگی های تابعی** باشد.

اگر تجزیه به گونه ای باشد که توسط وابستگی های تابعی دو جدول تجزیه شده $R1$ و $R2$ یعنی F' ، بتوانیم **همه** وابستگی های تابعی جدول اولیه یعنی F را استنتاج کنیم $F^+ = F'^+$ ، آنگاه شرط کافی محقق شده است. در **تجزیه مطلوب**، وابستگی های تابعی موجود در **جدول اولیه** نباید از دست بروند. به عبارت دیگر تمام FDهای موجود در جدول R یا باید در مجموعه FDهای $R1$ و $R2$ وجود داشته باشد، و یا از مجموعه FDهای موجود در $R1$ و $R2$ **مجدد** قابل استنتاج باشد.

$$F'^+ \equiv \{F_{\Pi_{<H1>}(R)} \cup F_{\Pi_{<H2>}(R)}\}^+$$

$$F^+ = F'^+$$

توجه مهم: در برخی متون کتب مرجع کلاسیک اصطلاح **بی حشو** مطرح نشده است و به همان **بی کاست (Non-Loss (LossLess))** بودن اکتفا شده است، در این مراجع، بی کاست بودن شامل بی حشو بودن هم هست، در واقع بی کاست و بی حشو هر دو باهم در اصطلاح بی کاست خلاصه شده است. البته در برخی منابع مرجع دیگر، **بی کاست (Non-Loss (LossLess))** بودن به معنی **بی کاست، بی حشو و حافظ وابستگی های تابعی** هر سه باهم است. همین موضوع باعث شده است سوالات ناواضح طراحان کنکور در این مبحث در **کلید نهایی** سازمان سنجش، **حذف** شود.

نکته مهم: بین شرط **لازم** و **کافی** تجزیه مطلوب، شرط لازم از اهمیت و اولویت بالاتری برخوردار است، **چرا؟** چون شرط لازم، نیاز اولیه است و شرط کافی، نیاز ثانویه. به عبارت دیگر در **طراحی بانک اطلاعات** رعایت شرط لازم (Loss-Less Join) اولویت بالاتر دارد نسبت به شرط کافی (حفظ وابستگی ها تابعی). بعد از اینکه **طراحی بانک اطلاعات** تمام شد، حال نوبت به رعایت قوانین جامعیت داخلی (مثل جامعیت درون رابطه ای) (داشتن حداقل یک کلید کاندید (اصلی)) و خارجی (قوانین کسب و کار) برای **جدول ها** می رسد. این آسیاب به نوبت است.

نرمال فرم دوم (Second Normal Form-2NF)

سطح دوم نرمال سازی، تبدیل جداول به نرمال فرم دوم است.

شرط قرار داشتن یک جدول در نرمال فرم دوم به صورت زیر است:

شرط لازم: جدول باید در نرمال فرم اول باشد.

شرط کافی: جدول باید فاقد وابستگی تابعی بخشی (غیر کامل) باشد.

وابستگی تابعی بخشی (PFD: Partial Functional Dependency): وابستگی یک مولفه

غیر کلیدی، به بخشی (جزئی یا عضوی) از کلید کاندید را وابستگی تابعی بخشی می نامند.

وابستگی تابعی بخشی (وابستگی بد) $\Rightarrow$ غیر کلید $X \rightarrow Y$ عضو کلید کاندید

مولفه غیر کلیدی: هر صفتی که عضو هیچ کلید کاندیدی نباشد مولفه غیر کلیدی است.

مولفه کلیدی (صفت عمده یا جزء یا عضو کلید کاندید): هر صفتی که عضو حداقل یک کلید

کاندید باشد، مولفه کلیدی است. صفتی که بخشی از کلید کاندید است تحت هیچ شرایطی به آن،

صفت غیر کلیدی گفته نمی شود، بلکه به آن همان مولفه کلیدی گفته می شود.

نکته مهم: اگر جدولی در نرمال فرم اول باشد و در آن تمام کلیدهای کاندید، تک خصیصه ای

باشند، آن جدول به طور خودکار فاقد وابستگی بخشی است و بنابراین در نرمال فرم دوم قرار

دارد. دقت کنید که شرط ایجاد وابستگی بخشی وجود کلید کاندید ترکیبی از چند خصیصه است،

بنابراین کلید کاندید تک خصیصه ای شرط ایجاد وابستگی بخشی را ندارد.

حذف وابستگی تابعی بخشی با تئوری هیث (Heath)

جهت حذف وابستگی تابعی بخشی و به تبع قرارگیری جدول در نرمال فرم دوم، تئوری هیث

مورد استفاده قرار می گیرد. در این تئوری، با تجزیه جدول به روش CUT کردن و برداشتن دُم

(دترمینان) وابستگی تابعی بخشی (بعلاوه کل ملحقات چسبیده به وابستگی بخشی) از روی کلید

کاندید و قرار دادن آن در یک جدول مستقل جدید، وابستگی تابعی بخشی، از جدول R حذف

می شود و مشکل درمان می شود. با حذف وابستگی بخشی، از روی جدول اولیه R، افزونگی

محتوایی ناشی از آن هم از بین می رود.

توجه مهم: تئوری هیث، به طور خودکار شروط تجزیه مطلوب (شرط لازم و کافی) را برقرار

می کند، یعنی هم حافظ تمام سطرها و ستونهای جدول اولیه و هم حافظ تمام FDهای جدول

اولیه است.

شرط لازم: INF باشد.

2NF

شرط کافی: پاس کردن سطح فعلی.

تجزیه مطلوب

شرط لازم: حافظ جدول اولیه باشد.

(بی کاست و بی حشو)

شرط کافی: حافظ FDها باشد.

فرآیند نرمال سازی از سه قدم زیر تشکیل شده است:

۱- کشف کلید کاندید

۲- کشف وابستگی های بد و خوب

۳- حذف وابستگی های بد

مثال: رابطه $R(A, B, C)$ با مجموعه وابستگی های تابعی (FD) زیر مفروض است:

$$F = \{ A \rightarrow C \}$$

وابستگی تابعی بخشی

غیر کلید	عضو کلید کاندید	غیر کلید
C	A	B
201	1	101
201	1	102
201	3	102
204	4	103

R

حاصل تجزیه 2NF جدول R به چه صورت است؟

۱- کشف کلید کاندید

با توجه به وابستگی های مطرح شده داریم:

$$F = \{ A \rightarrow C \}$$

اجتماع تمام خصیصه های سمت راست وابستگی های غیر بدیهی - تمام خصیصه های جدول = عضو کلید کاندید

$$AB - C = AB$$

بنابر رابطه فوق صفات AB حتماً باید عضو کلید کاندید باشد. بستر صفات AB به صورت زیر است:

$$\{AB\}^+ = \{AB\} \Rightarrow$$

$$A \rightarrow C \Rightarrow \{A, B, C\}$$

$$\{AB\}^+ = \{A, B, C\}$$

براساس بستر فوق، صفات AB، همه ستون های جدول R را تولید می کند، پس مطابق قانون دوم صفات AB تنها کلید کاندید جدول R است. و تحت هیچ شرایطی کلید کاندید دیگری ندارد.

۲ - کشف وابستگی های بد و خوب

وابستگی تابعی بخشی

	عضو کلید کاندید	عضو کلید کاندید	غیر کلید
	A	B	C
	1	101	201
	1	102	201
	3	102	201
	4	103	204

R

وابستگی تابعی بخشی (غیر کامل): **وابستگی بد** $\Rightarrow$ غیر کلید $A \rightarrow C$ عضو کلید کاندید

وابستگی تابعی کامل: **وابستگی خوب** $\Rightarrow$ غیر کلید $AB \rightarrow C$ کلید کاندید

توجه: وابستگی تابعی بد به طور ذاتی **مجاز** تکرار در سمت چپ دارد، اما وابستگی تابعی خوب به طور ذاتی **مجاز** تکرار در سمت چپ ندارد.

۳ - حذف وابستگی های بد (توسط تجزیه (ستز))

با **تجزیه جدول** به روش CUT کردن، وابستگی بخشی $A \rightarrow C$ از جدول R، **حذف** می شود و مشکل **درمان** می شود. به صورت زیر:

وابستگی تابعی کامل

تمام کلید	
A	B
1	101
1	102
3	102
4	103

R1

A	C
1	201
3	201
4	204

R2

توجه: وابستگی بخشی (بد) $A \rightarrow C$ از جدول R **حذف** شد و به وابستگی کامل (خوب) $A \rightarrow C$ در جدول R2 تبدیل شد.

توجه: تجزیه در نرمال سازی یعنی **دوری و دوستی**، دور هستن یعنی دو جدول شدن، اما دوست

هم هستن، یعنی توسط تعریف **کلید خارجی**، الحاق پذیر هم هستن و توسط عمل الحاق طبیعی همان **جدول اولیه** ایجاد می گردد.

در فرآیند نرمال سازی شروط **تجزیه مطلوب** به صورت زیر است:

۱- **شرط لازم:** (تجزیه، حافظ تمام ستونها و سطرها **جدول اولیه** باشد.)

در مثال قبل، دو جدول R1 و R2 الحاق پذیر هستند، همچنین **صفت مشترک A** در جدول R2 **کلید کاندید** و در جدول R1 **کلید خارجی** است، بنابراین شرط لازم **بی کاست** و **بی حشو** به طور هم زمان محقق شده است. خروجی $R_1 \bowtie R_2$ به صورت زیر است:

کلید خارجی		⋈	کلید کاندید		=			
A	B		A	C		A	B	C
1	101		1	201		1	101	201
1	102		3	201		1	102	201
3	102		4	204		3	102	201
4	103			R2		4	103	204
R1						R		

واضح است که $R = R1 \text{ JOIN } R2$ برقرار است و جدول اولیه **دقیقا** به دست می آید، یعنی الحاق (پیوند) بازسازنده (بازگشت پذیر) است.

نکته بسیار مهم: تجزیه جدول به روش **CUT کردن** باعث می شود، تجزیه به طور خودکار **بی کاست** و **بی حشو** شود.

ذهنیت جلسه کنکور: چیز بد (وابستگی بخشی) رو دیدی چیکارش باید کرد؟ هیچی دُمش رو بگیرى بندازیش بیرون.

۲- **شرط کافی:** (تجزیه، حافظ تمام **FDهای جدول اولیه** باشد.)

همه وابستگی های موجود در جدول اولیه R توسط جداول تجزیه شده R1 و R2 قابل دستیابی است، بنابراین شرط کافی محقق شده است.

$$F^{++} \equiv \{F_{\Pi_{<H1>}(R)} \cup F_{\Pi_{<H2>}(R)}\}^+$$

$$F^+ = F^{++}$$

$$F^+ = \{A \rightarrow C\}$$

وابستگی $A \rightarrow C$ در جدول R موجود است، از جدول تجزیه شده R2 هم در دسترس است.

شروط قرار داشتن یک جدول در نرمال فرم دوم به صورت زیر است:

نتیجه: دو جدول R1 و R2 در نرمال فرم دوم قرار دارند.

باشد، آنگاه تجزیه R به دو رابطه R1 و R2 **تجزیه مطلوب** است، به صورت زیر:

$$A \rightarrow B \quad : R1(A, B)$$

یک جدول می‌تواند **چندین تجزیه** داشته باشد، بدین صورت که **ترکیبی** از ستون‌ها به طرق مختلف **ستون مشترک** باشد، که در بین آنها یکی **تجزیه مطلوب** و مابقی **تجزیه نامطلوب** خواهد بود، به صورت زیر:

B	A	A	C	A	B	B	C	A	C	C	B
R1		R2		R1		R2		R1		R2	
تجزیه مطلوب با ستون مشترک A				تجزیه نامطلوب با ستون مشترک B				تجزیه نامطلوب با ستون مشترک C			

تجزیه نامطلوب بدون ستون مشترک

تجزیه نامطلوب بدون ستون مشترک

تجزیه نامطلوب بدون ستون مشترک

برای نمونه یک تجزیه نامطلوب با ستون مشترک C به صورت زیر است:

A	C	⋈	B	C	=	A	B	C
1	201		101	201		1	101	201
3	201		102	201		1	102	201
4	204		103	204		3	101	201
R1			R1			3	102	201
						4	103	204
						R		

واضح است که $R = R1 \text{ JOIN } R2$ برقرار نیست (سطر سوم اضافه است) و جدول اولیه دقیقاً به دست نمی آید، یعنی الحاق (پیوند) بازسازنده (بازگشت پذیر) نیست. چون ستون مشترک C در حداقل یکی از دو جدول کلید کاندید نیست.

نرمال فرم سوم (Third Normal Form-3NF)

سطح سوم نرمال سازی، تبدیل جداول به نرمال فرم سوم است.

شروط قرار داشتن یک جدول در نرمال فرم سوم به صورت زیر است:

شرط لازم: جدول باید در نرمال فرم دوم باشد.

شرط کافی: جدول باید فاقد وابستگی تابعی انتقالی (تراگذری یا باواسطه) باشد.

وابستگی تابعی انتقالی (TFD: Transitive Functional Dependency): وابستگی یک مولفه

غیرکلیدی، به یک مولفه غیرکلیدی دیگر را وابستگی تابعی انتقالی می نامند.

وابستگی تابعی انتقالی (وابستگی بد) $\Rightarrow$ غیرکلید $X \rightarrow Y$ غیرکلید

مولفه غیرکلیدی: هر صفتی که عضو هیچ کلید کاندیدی نباشد مولفه غیرکلیدی است.

حذف وابستگی تابعی انتقالی با تئوری ریسانن (Rissanen)

جهت حذف وابستگی تابعی انتقالی و به تبع قرارگیری جدول در نرمال فرم سوم، تئوری ریسانن مورد استفاده قرار می گیرد. در این تئوری، با تجزیه جدول به روش CUT کردن و برداشتن دُم (دترمینان) وابستگی تابعی انتقالی (بعلاوه ملحقات چسبیده به وابستگی انتقالی) از روی جدول پایه R و قرار دادن آن در یک جدول مستقل جدید، وابستگی تابعی انتقالی از جدول R حذف می شود و مشکل درمان می شود. با حذف وابستگی انتقالی از روی جدول پایه R، افزونگی های ناشی از آن هم از بین می رود.

توجه مهم: تئوری ریسانن، به طور خودکار شروط تجزیه مطلوب (شرط لازم و کافی) را برقرار

می‌کند، یعنی هم حافظ جدول پایه و هم حافظ FD ها است.

توجه مهم: جداولی که توسط تئوری ریسائن، تجزیه می‌شوند، به جداول مستقل از هم موسوم هستند.

نکته مهم: دقت کنید، مدل ریاضی تجزیه، یک **درخت تجزیه دودویی** است، یعنی در هر مرحله، تجزیه دودویی است. برای مثال اگر یک جدول شامل چند وابستگی تابعی انتقالی باشد، در هر مرحله از درخت تجزیه، یکی از وابستگی‌های تابعی انتقالی حذف می‌شود و همه باهم در یک مرحله حذف نمی‌شود.

فرآیند نرمال‌سازی از **سه قدم** زیر تشکیل شده است:

۱- کشف کلید کاندید

۲- کشف وابستگی‌های بد و خوب

۳- حذف وابستگی‌های بد

جهت دانلود ادامه جزوه به لینک زیر در سایت موسسه بابان مراجعه نمایید:

BabanAcademy.com/database

مثال: رابطه $R(A, B, C)$ با مجموعه وابستگی‌های تابعی (FD) زیر مفروض است:

$$F = \{A \rightarrow B, B \rightarrow C\}$$

وابستگی انتقالی وابستگی کامل

A	B	C
1	101	202
2	102	203
3	102	203
4	103	202

R