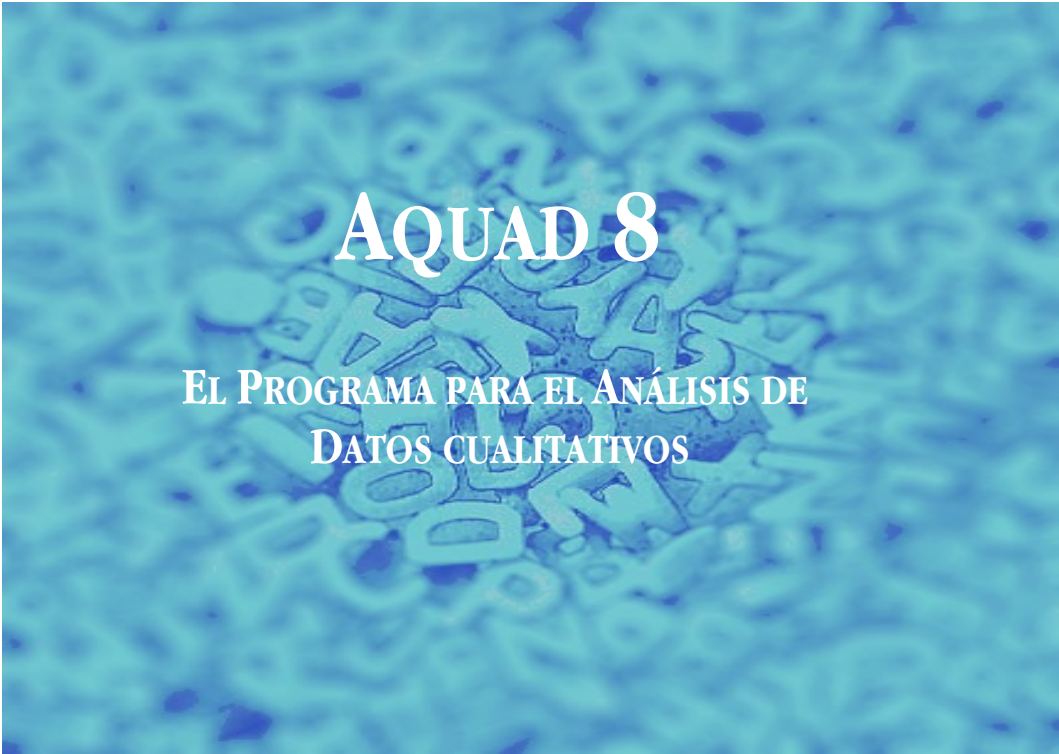




AQUAD 8

EL PROGRAMA PARA EL ANÁLISIS DE
DATOS CUALITATIVOS

Günter L. Huber
Leo Gürtler

Copyright

Todos los derechos reservados. Ninguna parte de este manual puede ser reproducida o transmitida en cualquier forma o por cualquier medio, ya sea electrónico, mecánico, de fotocopia, de grabación o de otro tipo, sin la autorización expresa por escrito de los autores. Lo mismo ocurre con los derechos de reproducción pública y de traducción a otras lenguas.

El propio AQUAD 8 es un software libre. Puede redistribuirlo y/o modificarlo bajo los términos de la Licencia Pública General GNU publicada por la Fundación del Software Libre, ya sea la versión 3 de la Licencia, o (a su elección) cualquier versión posterior.

AQUAD 8 está programado con LAZARUS 3.4 / FPC versión 3.2.2 utilizando el componente «paslibvlc» de Robert Jedrzejczyk (<https://prog.olsztyn.pl/paslibvlc/>) y se compila como versión de 32 bits. Para analizar grabaciones de audio o vídeo (con aquad8_c_audio.exe o aquad8_c_video.exe), AQUAD 8 utiliza la versión de 32 bits del reproductor multimedia «VLC» (<https://www.videolan.org>). A partir de la versión AQUAD 8.6, AQUAD 8 muestra, copia, imprime, etc. los resultados con el programa «notepad++ .exe» (<https://notepad-plus-plus.org/downloads/>), ya que «notepad.exe» ya no está incluido en las versiones actuales de «Microsoft Windows».

Es posible que tenga que descargar «VLC» (versión de 32 bits!) y «notepad++ .exe» e instalarlos en su ordenador.

Limitación de la garantía

Los autores no ofrecen ningún tipo de garantía con respecto al programa AQUAD 8 descrito en este manual y, por lo tanto, no debe derivarse ninguna responsabilidad del uso del paquete de programas o de cualquier parte del mismo.

En este manual se utilizan varias marcas registradas:

Windows, Word y Excel son marcas comerciales de Microsoft Corporation.

WordPerfect es una marca comercial de Corel Corporation.

R es una marca comercial de la R Foundation

VLC es una marca comercial de Videolan

3ª edición 2021

Autores: Günter L. Huber y Leo Gürtler

Günter Huber, Viktor-Renner-Str. 39, 72074 Tübingen, Alemania

teléfono 07071-885147

Correo electrónico: info@aquad.de

Contenido

Introducción	7
Capítulo 1: Instalación del programa AQUAD	9
Capítulo 2: Preparación y almacenamiento de datos	11
2.1 Trabajar con archivos de texto	11
2.1.1 Formatear los archivos de texto	11
2.1.2 Nombra tus archivos con sentido	12
2.1.3 Convertir archivos de texto a formato de texto (*.txt)	13
2.2 Trabajar con archivos de audio	14
2.3 Trabajar con archivos de vídeo	15
2.4 Trabajar con archivos gráficos	15
Capítulo 3: La estructura del programa AQUAD	16
3.1 Paradigmas de análisis de datos	16
3.2 Módulos y funciones del análisis de datos	17
3.2.1 Proyecto	18
3.2.2 Métodos de análisis	18
3.2.3 Búsqueda	20
3.2.4 Tratar códigos	21
3.2.5 Tablas	22
3.2.6 Vínculos	22
3.2.7 QCA/implicantes	23
3.2.8 Anotaciones	24
3.2.9 Herramientas	25
Capítulo 4: El proceso de análisis de contenido cualitativo	25
4.1 La reducción de los datos cualitativos	27
4.2 Esquema del proceso de análisis de datos	28
4.3 Elaboración de códigos para unidades de significado	32
4.3.1 Desarrollo de códigos categóricos	33
<i>Uso de sistemas de categorización predeterminados</i>	33
<i>Categorización basada en hipótesis</i>	34
<i>Categorización para la construcción de la teoría</i>	35
4.3.2 Desarrollo de códigos secuenciales	36
<i>Búsqueda de secuencias simples</i>	36
<i>Búsqueda de patrones de secuencias complejas</i>	37
4.3.3 Elaboración de códigos temáticos	37
4.4 Reconstrucción de los sistemas de significado	37
4.4.1 Reconstrucción de secuencias simples de significado	38
4.4.2 Reconstrucción de sistemas de significado condicional	39
4.5 Comparación de los sistemas de significado	39
Capítulo 5: El análisis de textos como análisis de secuencias	41
5.1 Preparación de los textos	41

5.2 Generación de hipótesis	43
5.2.1 Antecedentes teóricos	43
5.2.2 Principios de la generación de hipótesis	44
5.2.3 Condensar las hipótesis en lecturas del texto	47
5.3 Comprobación de hipótesis	48
5.4 Generación y comprobación de hipótesis con AQUAD	48
5.4.1 Generación de hipótesis	49
5.4.2 Prueba de hipótesis	51
 Capítulo 6: Preguntas especiales de codificación	 52
6.1 Creación de un proyecto	52
6.2 ¿Cómo se codifica?	52
6.2.1 Tipos de códigos	52
6.2.2 Selección de archivos de texto para la codificación	54
6.2.3 Reglas para la inserción de códigos	54
6.2.4 Cómo eliminar o sustituir los códigos	56
6.2.5 Cómo obtener una visión general de los códigos	58
6.2.6 Cómo añadir comentarios a los códigos	59
6.2.7 Cuando se necesita un poco de ayuda	60
6.2.8 Por qué no dejar que Aquad se encargue de la búsqueda:	
codificación semiautomática	61
6.2.9 Si desea imprimir archivos de código	61
6.2.10 ¿Qué es el registro de códigos y por qué querría eliminarlo?	62
6.2.11 Codificación de archivos de imagen	63
6.2.12 Codificación de archivos de audio	64
<i>Trabajar con el reproductor multimedia</i>	64
<i>Estrategia general de codificación de archivos de audio</i>	66
<i>Transcripción de segmentos importantes del archivo</i>	66
6.2.13 Codificación de archivos de vídeo	67
6.3 ¿Cómo editar el texto y la codificación?	67
6.3.1 Cómo editar textos	67
6.3.2 Cómo editar las codificaciones	67
6.4 ¿Para qué sirven los metacódigos?	68
6.4.1 Modificación de archivos: Añadir metacódigos	69
6.4.2 Cambio de archivos: Sustituir los códigos por metacódigos	69
6.4.3 Restaurar códigos anteriores	70
6.5 ¿Cómo añadir varios códigos?	70
6.6 ¿Cómo utilizar las palabras clave?	71
 Capítulo 7: Cómo trabajar con códigos	 72
7.1 Cómo encontrar los códigos	72
7.1.1 Requisito previo: crear un catálogo de códigos	72
7.1.2 Cómo encontrar segmentos de texto codificados y códigos específicos	73
7.1.3 Cómo buscar estructuras de código	74
<i>Códigos anidados</i>	74
<i>Codificaciones superpuestas</i>	75
<i>Codificaciones múltiples</i>	76

<i>Secuencias de codificación</i>	77
7.1.4 Excursus: Distancia entre segmentos de archivos	78
7.1.5 Códigos NO utilizados	79
7.1.6 Contar códigos e introducir frecuencias en las tablas	79
7.2 Cómo encontrar correlaciones entre códigos	79
7.2.1 Relaciones en el contexto de las categorías incondicionales: Búsqueda de códigos	81
7.2.2 Relaciones en el contexto de las categorías condicionales: Análisis de tablas o matrices	81
<i>Cómo construir una tabla</i>	83
<i>Cómo editar las tablas</i>	85
<i>Cómo realizar un análisis de la tabla</i>	85
7.3 Relaciones en forma de secuencias de codificación simples: Comprobación de vínculos	86
7.3.1 ¿Qué estructuras de vínculos proporciona AQUAD? <i>Secuencias de codificación generales</i>	87
<i>Comparación de codificaciones entre dos hablantes</i>	90
7.3.2 Cómo construir vínculos	92
Capítulo 8: Análisis de contenido cuantitativo	95
8.1 Contar palabras	95
8.2 Contar sufijos	98
8.3 Cómo excluir partes del texto del análisis de palabras	99
Capítulo 9: Trabajar con memos	100
9.1 ¿Para qué sirven los memos?	100
9.2 ¿Cómo se escriben las notas en AQUAD? 9.2.1 Escribir notas mientras se codifica	100
9.2.2 Escribir memos desde el menú principal	102
9.3 ¿Cómo encuentro mis memos? 9.3.1 Búsqueda de memos durante la codificación	102
9.3.2 Búsqueda de memos desde el menú principal	103
Capítulo 10: Análisis de implicantes (análisis comparativo cualitativo)	104
10.1 Escribir una tabla de datos	104
10.2 Convertir datos en valores de verdad	106
10.3 Búsqueda de implicantes en los criterios "positivos" y "negativos"	108
10.4 Para qué más se pueden utilizar los implicantes ...	110
10.5 Función de los implicados en el proceso de construcción de la teoría	113
Capítulo 11: Combinación de análisis cualitativos y cuantitativos	114
11.1 El debate acerca los métodos de investigación cualitativos frente a los cuantitativos	114
11.2 "Métodos mixtos": ¿Palabra de moda o estrategia de investigación? 11.2.1 Combinación implícita de métodos	118
11.2.2 Combinación explícita de métodos cualitativas/cuantitativas	118
11.2.3 Combinación de métodos en el diseño de la investigación	121
<i>(1) Macrosecuencias de la combinación de métodos</i>	121

(2) <i>Aplicación simultánea de métodos cualitativos y cuantitativos</i>	123
(3) <i>Microsecuencias de la combinación de métodos</i>	124
11.3 Frecuencias de palabras	124
11.4 Frecuencias del código	124
11.5 Combinaciones de métodos con el ejemplo de un estudio sobre el humor	125
11.5.1 Análisis de varianza univariante de las frecuencias de los códigos para la generación de hipótesis de vinculación	126
11.5.2 Formación de tipos mediante el análisis de implicantes sobre una base cuantitativa	131
Capítulo 12: Conversión de códigos del formato AQUAD 7 al AQUAD 8	133
Capítulo 13: Análisis exploratorio estadístico	134
13.1 Modificación de tablas de datos	134
13.2 Estadística exploratoria descriptiva	134
13.3 Análisis cualitativo comparativo (Minimización booleana)	136
Referencias	138
Sobre los autores	144

Introducción

La primera versión de AQUAD (Análisis de Datos Cualitativos) se desarrolló en 1987 para permitir que un proyecto de investigación se completara dentro de un ajustado periodo de financiación. En aquella época, ya existían algunas herramientas informáticas para el análisis de datos cualitativos. Los más sencillos utilizaban las funciones de búsqueda de los programas de texto o de bases de datos disponibles. Algunos programas se escribieron específicamente para los requisitos del análisis cualitativo, pero normalmente no ofrecían más que simples funciones de recuento y búsqueda. En cambio, el salto cualitativo lo dio un programa estadounidense, el programa informático Qualog de Anne Shelly y Ernest Sibert para grandes ordenadores. Utilizaba las avanzadas posibilidades de la llamada "programación lógica" y fue la inspiración y el modelo de AQUAD.

En particular, AQUAD se desarrolló para apoyar el descubrimiento del significado en el material de datos, como se describe, por ejemplo, en el enfoque metodológico de la "teoría fundamentada" (Glaser & Strauss, 1967). Además, se incorporaron a AQUAD las ideas metodológicas de Miles y Huberman (1984, 1994) sobre el análisis tabular o matricial y se integró el enfoque del análisis comparativo cualitativo de Ragin (1987) para comparar configuraciones de significado en diferentes conjuntos de datos. Desde la versión 7, también se ha incorporado un módulo de interpretación hermenéutica objetiva de textos según Oevermann (1981), siguiendo las sugerencias de Wernet (2009) para el análisis secuencial. Por supuesto, en AQUAD también existen funciones sencillas, como la gestión de archivos, la búsqueda de codificaciones, el recuento de palabras, etc.

Para simplificar el software y adaptarlo mejor a las necesidades del usuario, AQUAD se ha dividido en cuatro subprogramas separados en la versión 8:

- Los archivos de texto con formato "*.txt" se analizan cualitativamente con mayor frecuencia. Para ello está disponible "aquad8_c_texto.exe".
- Incluso sin una transcripción exhaustiva, las grabaciones de audio de las entrevistas en muchos formatos (por ejemplo, "*.wav", "*.mp3", "*.aac", etc.) pueden analizarse directamente con la versión "aquad8_c_audio.exe". Por supuesto, se omiten opciones como la búsqueda de palabras clave o la salida escrita de pasajes de texto típicos. La versión de 32 bits del reproductor multimedia VLC debe estar disponible en su ordenador (www.videolan.org).
- Las grabaciones de vídeo (en numerosos formatos, por ejemplo, "*.avi", "*.mp4", "*.mov", etc.) también pueden analizarse sin necesidad de transcripción con la versión "aquad8_c_video.exe". La versión de 32 bits del reproductor multimedia VLC debe estar disponible en su ordenador (www.videolan.org).
- Las fotos, los dibujos escaneados y otros datos gráficos proporcionados en archivos de formato "*.jpg" o "*.png" pueden analizarse con la versión "aquad8_c_grafic.exe".

La lógica de funcionamiento del programa y el tratamiento interno de las codificaciones sólo difieren ligeramente para los distintos formatos de archivo de AQUAD, por lo que la familiarización es fácil. Para los usuarios de la versión 7, no hay cambios fundamentales; los archivos existentes se pueden convertir para la versión 8. Al codificar los textos, en la versión 8 sólo se tienen en cuenta los números de línea de los segmentos de texto seleccionados para su localización en el texto, pero ya no su comienzo exacto ni su longitud (número de caracteres) en el archivo de texto.

Tras la conversión, los archivos de código existentes de la versión 7 limitan la localización de los segmentos de texto a los números de línea.

Capítulo 1: Instalación del programa AQUAD

En nuestra página web www.aquad.de encontrará en la sección "Download" en el apartado "AQUAD 8" los siguientes módulos de instalación del paquete de programas:

- aquad8_c_texto_setup.exe: Módulo para el análisis cualitativo y cuantitativo de textos según el paradigma de codificación y el paradigma de reconstrucción (véase el capítulo 3).
- aquad8_c_audio_setup.exe: Módulo para el análisis cualitativo de grabaciones sonoras
- aquad8_c_video_setup.exe Módulo para el análisis cualitativo de grabaciones de vídeo
- aquad8_c_grafic_setup.exe Módulo para el análisis de material de datos gráficos (fotos, dibujos, grabaciones manuscritas, etc.)

Aquad 8 utiliza algoritmos de 32 bits de Robert Jedrzejczyk (<https://prog.olsztyn.pl/paslibvlc/>) para la reproducción de audio y vídeo con el reproductor multimedia gratuito VLC.

Las versiones de audio y vídeo de Aquad 8 funcionan sin problemas si instala la **versión de 32 bits de VLC** en su ordenador.

Las funciones de procesamiento de texto de los resultados son proporcionadas por el software gratuito «notepad + + .exe». Si desea imprimir, cortar, copiar, pegar, guardar en directorios distintos a los de Aquad, etc., utilice las funciones del menú de la ventana correspondiente de «notepad + + .exe».

Como copia de seguridad, el resultado del último análisis se guarda en el directorio ..\res con el título «work.} } }» y se sobrescribe en el siguiente análisis. Por supuesto, «work.} } }» se puede cargar con otros programas para cualquier uso.

Aquad muestra todos los resultados de las evaluaciones de frecuencia en forma de listas en «notepad + + .exe» y, además, los guarda como tablas en formato «*.csv» para su posterior análisis exploratorio de datos o para importarlos a programas estadísticos.

Durante la instalación, Aquad crea automáticamente los siguientes directorios en su ordenador. Para simplificar el ejemplo, suponemos que instala el módulo de análisis de texto y que acepta la sugerencia «C:\AQUAD8_d_text» de la rutina de instalación para el directorio raíz de todos los archivos relacionados con AQUAD (para archivos de audio, vídeo y gráficos se aplica lo mismo):

Contenido del directorio

<i>Senda</i>	<i>Contenido</i>
C:\AQUAD8_c_texto	Directorio base o raíz; Los textos de muestra suministrados se instalan automáticamente aquí. Los textos (*.txt) de su proyecto también deben guardarse en este directorio.
C:\AQUAD8_c_*\cod * = texto, audio, grafic, video	Los códigos que se crean durante su trabajo de interpretación (*.aco); Aquí también encontrará los parámetros de los proyectos y los directorios de códigos, palabras clave, etc., que es probable que guarde para su posterior análisis.
C:\AQUAD8_c_*\cod_s	Este directorio es para hacer una copia de seguridad de sus códigos, especialmente para poder restaurar los códigos anteriores en caso de combinar diferentes códigos que son similares en significado y reemplazarlos con el nuevo metacódigo(s).
C:\AQUAD8_c_*\lit	Este manual. Aquí también puedes almacenar publicaciones, enlaces, etc. que son importantes para tu trabajo en el proyecto actual.
C:\AQUAD8_c\mco	Archivos de metacódigos (resúmenes de códigos de significado).
C:\AQUAD8_c_*\prg	El propio programa, el archivo de ayuda y los parámetros del proyecto actual.
C:\AQUAD8_c_*\res	Los resultados de todas las evaluaciones (resultados) como archivos txt y/o csv.

Capítulo 2: Preparación y almacenamiento de datos

Con el módulo apropiado (ver arriba) de AQUAD 8,

archivos de texto en formato *.txt,

archivos de audio en todos los formatos de sonido disponibles en el reproductor multimedia VLC (por ejemplo, *.mp3, *.wav, *.aac, etc.),

archivos de vídeo en todos los formatos de vídeo disponibles en el reproductor multimedia VLC (por ejemplo, *.mp4, *.avi, *.mov, etc.) o

archivos gráficos en formato *.jpg o *.png

pueden ser analizado. Estos archivos deben guardarse en el directorio raíz de AQUAD (en el ejemplo de los análisis de texto del capítulo 1: C:\Aquad_8d), ya que el programa busca allí por defecto los datos a analizar. Por supuesto, puedes guardar tus archivos originales en otro lugar, incluso en otro soporte de datos, por seguridad, pero una copia de trabajo debe estar disponible para AQUAD en el directorio raíz. Para la preparación de los archivos de texto para su análisis, he aquí algunos consejos que han resultado útiles a muchos usuarios de AQUAD:

2.1 Trabajar con archivos de texto

Antes de definir su propio proyecto, debe preparar sus textos para que AQUAD pueda acceder a ellos. En concreto, debe

- formatear sus documentos de texto en su programa de tratamiento de textos,
- convertirlos a formato "txt"
- y luego copiarlos en el directorio raíz de AQUAD (ver instalación: "C:\Aquad8_c_texto" u otros módulos).

2.1.1 Formatear los archivos de texto

Puede utilizar cualquier programa de tratamiento de textos para introducir sus datos. Sin embargo, antes de poder trabajar de forma significativa con los textos dentro de AQUAD, hay que preparar cada archivo para ello. En primer lugar, reformatea los textos para que las líneas no contengan más de unos 60 caracteres, un valor que ha resultado útil en la práctica. Las líneas de este texto constan de unos 90 caracteres cada una, por lo que deben acortarse en un tercio aproximadamente y justificarse a la izquierda, así:

```
Puede utilizar cualquier programa de tratamiento de textos
para introducir sus datos. Sin embargo, antes de poder trabajar
de forma significativa con los textos dentro de AQUAD, hay que
preparar cada archivo para ello. En primer lugar, reformatea
```

El motivo es que AQUAD utiliza números de línea para marcar segmentos de texto significativos. Por lo tanto, las líneas individuales no deben ser demasiado largas para que los códigos puedan asignarse con precisión a las partes pertinentes del texto.

Este formato se realiza mejor con el programa de texto mientras se introducen los datos. Sin embargo, también puede volver a formatear como desee después (véase el siguiente párrafo). En la mayoría de los programas de tratamiento de textos, puede determinar la longitud de las líneas ajustando los márgenes en consecuencia (sin margen a la izquierda, a la derecha el margen debe ajustarse de modo que haya espacio para unos 60 caracteres).

Pero, ¿qué hacer si los textos ya han sido introducidos con el ordenador y formateados para "llenar la página"? No hay problema: con cualquier programa de texto, incluso con muchos editores de texto muy sencillos, puede reformatear posteriormente los textos según el principio que acabamos de describir. Muchos programas de texto suelen permitir especificar un formato de página (incluida la longitud de las líneas) al que se adaptan automáticamente todos los textos. Así, sólo tendrás que cargar tus textos uno tras otro, y volver a guardarlos. Consulte el manual de su programa de texto para más detalles.

2.1.2 Nombrar los archivos con sentido

Cuando guardas tus archivos de texto, en Aquad tiene sentido que los tres últimos caracteres del nombre sean dígitos. Como acuerdo para nombrar los archivos de texto, Aquad tiene lo siguiente:

- Los nombres de los archivos no tienen más de ca. 60 caracteres.
- Los tres últimos caracteres del nombre deben ser dígitos.
- Como extensiones, los programas de texto añaden ".txt" a los nombres de los archivos al convertirlos en archivos de texto plano; Aquad utiliza estas extensiones para identificar los archivos de texto.
- Para la anonimización, todos los archivos de texto de un proyecto reciben el mismo nombre pero diferentes dígitos finales. O de otra manera: los archivos de un proyecto se numeran consecutivamente (no necesariamente con números consecutivos); a cada archivo se le asigna un número de tres dígitos diferente.

Probablemente elija los nombres para que el contenido de los archivos sea reconocible para usted. Por ejemplo, si sus datos consisten en 24 transcripciones de entrevistas y ha elegido "entrevista" como nombre de su proyecto, podría nombrar sus archivos originales de la siguiente manera:

entrevista_001.txt, entrevista_002.txt, ..., entrevista_024.txt.

Es muy importante que los nombres de los archivos sean únicos, es decir, que se diferencien en números. Por supuesto, puede prescindir de los números y utilizar otros nombres, como "Meyer_Hegelschule.txt". Sin embargo, al hacerlo, se renuncia al anonimato de los entrevistados y se dificulta si se almacenan datos de varios proyectos en el mismo directorio raíz y se quiere hacer una selección específica de los mismos. Además, nunca utilices "trabajo" como nombre porque Aquad lo utiliza internamente.

Mientras Aquad trabaja con tus archivos, también crea archivos paralelos que contienen, por ejemplo, la información del código de cada archivo. El programa crea nuevos nombres de archivo a partir de los elementos de los nombres originales de forma que se pueda reconocer la conexión con los archivos originales. En nuestro ejemplo, los archivos de código tendrían los siguientes nombres:

entrevista_001.aco, entrevista_002.aco, ..., entrevista_024.aco.

Después de codificar, puedes ver estos nombres en el subdirectorio Aquad "..\cod", que se ha configurado automáticamente para almacenar los archivos de código. De este modo, al trabajar con Aquad se crean muchos archivos especiales diferentes y no hay que preocuparse por su denominación exacta. Por este motivo, no daremos aquí una lista exhaustiva.

Si sus archivos se guardan con nombres que se desvían de esta convención, puede renombrar estos archivos. Puede hacerlo dentro de un "gestor de archivos" como el "Explorador de Windows" de Microsoft, haciendo clic con el botón derecho del ratón en el nombre del archivo y seleccionando la función "Renombrar" en el menú emergente que aparece. También puede escribir "CMD" (por "comando") en la ventana de búsqueda situada en la parte inferior de la barra de tareas de Windows y, a continuación, hacer clic en la opción "Símbolo del sistema". Asegúrese de que se muestra el subdirectorio de sus archivos de texto (textos originales) y, a continuación, utilice el comando "REN" (para "renombrar") de la siguiente manera:

```
REN nombre_antigo nombre_nuevo
```

Supongamos que los textos de las entrevistas de nuestro ejemplo llevan originalmente el nombre de los entrevistados (lo que desaconsejamos por razones de protección de datos). Veamos el procedimiento para los dos primeros textos, que se guardan como wolfgang.txt y andrea.txt. Los renombramos así

```
REN wolfgang.txt inter_001.txt
REN andrea.txt inter_002.txt
```

y proceder en consecuencia con los otros 22 textos de nuestro ejemplo.

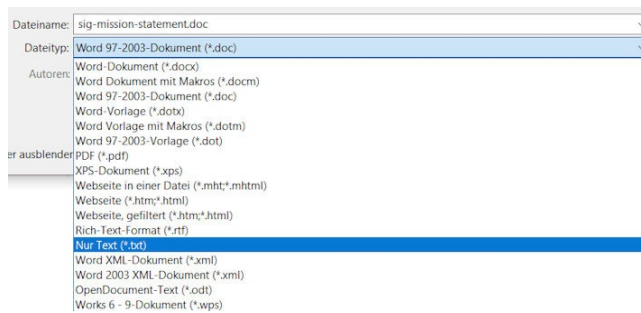
2.1.3 Convertir los textos en formato de texto ANSI (*.txt)

Todos los documentos de texto que quiera utilizar en AQUAD deben ser convertidos al formato ANSI (*.txt). "ANSI" es una abreviatura y una designación registrada del American National Standards Institute; el código ANSI es el que se utiliza habitualmente en los programas de Windows para representar caracteres e intercambiar información de texto entre diferentes programas.

Cuando escribe un texto, su procesador de textos añade al mismo, de forma invisible para Ud., códigos de tamaño de letra, color, fin de palabra y otras informaciones, como instrucciones al programa para que inicie una nueva línea, establezca márgenes, sangrías, inicie una nueva página, añada cursiva o negrita o color al texto, etc. También incluye instrucciones a la impresora sobre cuándo debe comenzar una nueva página y mucho más. Lamentablemente, estas instrucciones son específicas del programa y no pueden leerse en otros programas, a menos que éstos contengan una rutina de transformación, para la que, sin embargo, hay que pagar una licencia. Para exportar e importar archivos de texto sin problemas, hay que eliminar todos los caracteres que sólo tienen sentido en un programa concreto, es decir, los archivos deben estar normalizados. Para la importación en AQUAD, sus documentos de datos específicos deben convertirse en documentos estándar ANSI.

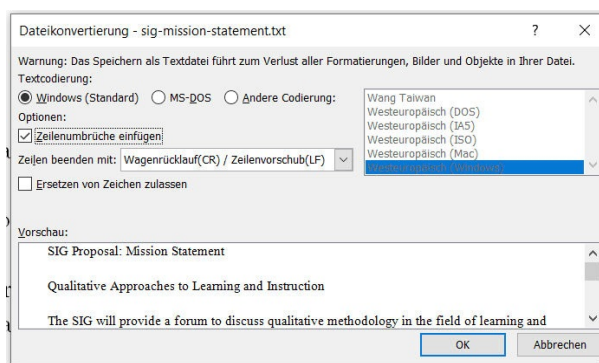
Por ello, los programas de texto suelen tener una función de ayuda que convierte sus documentos en textos ANSI (txt). Consulte el manual de su programa de tratamiento de textos en "texto ANSI", "formato txt", "importar/exportar", "entrada y salida de texto" "solo texto (txt)" o palabras clave similares.

Las siguientes dos capturas de pantalla muestran el procedimiento para trabajar con MS Word (versión en alemán):



Si se selecciona la opción "Guardar como" (Speichern unter), hay que cambiar el "Tipo de archivo" (Dateityp) y hacer clic en el formato "Solo texto (*.txt)" (Nur Text (*.txt)).

Después de hacer clic en el botón "Guardar" aparece esta ventanilla adicional:



En esta ventana dejamos la configuración por defecto "Windows" y el idioma a la derecha, pero hacemos clic en "Insertar saltos de línea" y comprobamos si está marcado que los finales de línea se ejecuten con los códigos estándar "CR/LF" (Retorno del carro/Avance de línea).

No es necesario que lea la siguiente sección si sólo quiere orientarse al principio. Sólo será importante cuando esté trabajando con Aquad y se dé cuenta de que necesita cambiar algo en uno u otro archivo.

Edición de archivos tras la conversión al formato de texto ANSI

Si encuentra errores en los datos originales, puede corregirlos. Sin embargo, hay una regla básica en el análisis cualitativo: los documentos de texto no deben modificarse durante el análisis. Por ello, compruebe los textos antes de iniciar un análisis cualitativo de los mismos. Para sus interpretaciones en forma de códigos, AQUAD almacena los números de línea de cada segmento de texto codificado. Si elimina o añade letras o palabras enteras, puede cambiar la estructura de las líneas del texto y, por tanto, la asignación de los códigos.

2.2 Trabajar con archivos de audio

Para utilizar los módulos de análisis de archivos de audio, es necesario instalar en el ordenador la **versión de 32 bits** del reproductor multimedia gratuito VLC (descarga: www.videolan.org).

Las recomendaciones para nombrar los archivos de texto (véase más arriba) también se aplican a los archivos de audio, en particular los nombres de los archivos no deben tener más de 60 caracteres.

Si quiere analizar archivos de audio que ya están disponibles, guárdalos en el directorio raíz de AQUAD de forma análoga al procedimiento de los archivos de texto. Si quiere grabar primero las entrevistas, por ejemplo, basta con que conectes un micrófono a la entrada correspondiente del portátil y utilice el dispositivo como grabadora con el software adecuado (recomendamos el programa gratuito Audacity, que se puede descargar de

<http://audacity.sourceforge.net/download>). Una máquina de dictado digital o una aplicación correspondiente en un smartphone también son adecuadas para la mayoría de los propósitos. De nuevo, debe copiar los archivos de audio de estos dispositivos en el directorio raíz de AQUAD.

2.3 Trabajar con archivos de vídeo

Para utilizar los módulos de análisis de las grabaciones de vídeo, es necesario instalar en el ordenador la versión de 32 bits del reproductor multimedia gratuito VLC (descarga: www.videolan.org).

Las recomendaciones para nombrar los archivos de texto (véase más arriba) también se aplican a los archivos de vídeo, en particular los nombres de archivo no deben tener más de 60 caracteres.

Si quiere analizar archivos de vídeo que ya están disponibles, guárdalos en el directorio raíz de AQUAD de la misma manera que los archivos de texto. Si quiere grabar en vídeo entrevistas o secuencias de interacción, puede utilizar videocámaras digitales o la función de vídeo de una cámara digital o un smartphone. De nuevo, debe copiar los archivos de vídeo de estos dispositivos en el directorio raíz de AQUAD.

2.4 Trabajar con archivos gráficos

No es necesaria una preparación previa para sus archivos de imagen. AQUAD 8 está preparado para la importación de imágenes en formato "*.jpg" y "*.png". Por lo tanto, con su programa de escaneo o de tratamiento de imágenes, guarde las imágenes destinadas al análisis en AQUAD, preferiblemente en el formato

.jpg (comprimible, pequeño y manejable).

Otros formatos como .tif o .bmp requieren mucha más memoria. Los archivos de imagen pueden abrirse directamente en AQUAD, como está acostumbrado a hacer con los textos. Las dimensiones de las imágenes se ajustan automática y proporcionalmente.

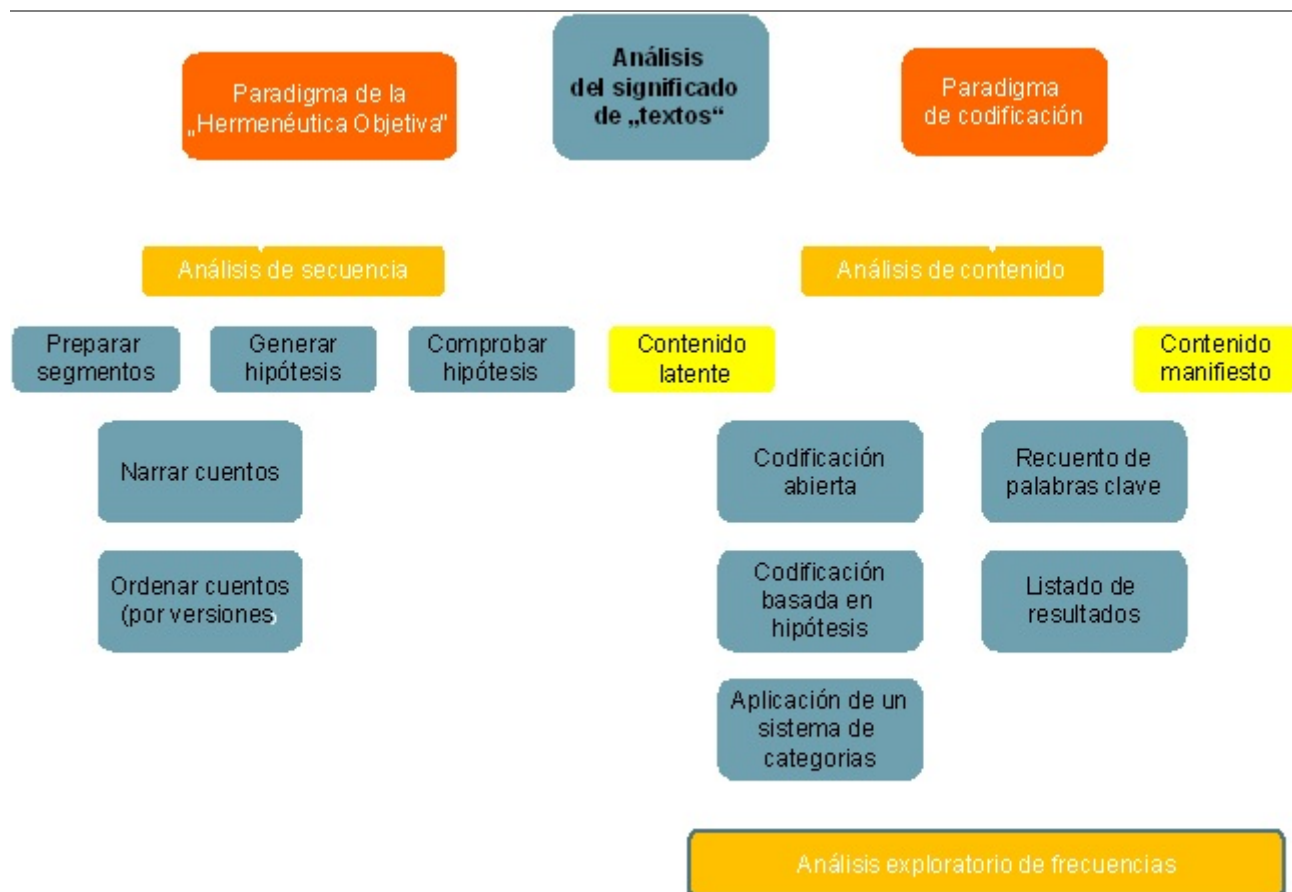
Dado que prácticamente todos los programas de gráficos y también muchos programas de visualización de imágenes ("visores") pueden convertir archivos de imagen de casi cualquier formato común a casi cualquier otro, nos pareció superfluo llenar AQUAD con una colección de rutinas de conversión. En Internet también se pueden encontrar programas de visualización y conversión para descargar, por ejemplo, en las páginas de muchas revistas de informática, en sitios de programas gratuitos y en servidores universitarios.

Tenemos buenas experiencias con el programa "IrfanView" de Irfan Skiljan. Para uso no comercial, por ejemplo en proyectos de investigación universitarios, el programa puede obtenerse gratuitamente en

<http://irfanview.tuwien.ac.at>

Capítulo 3: La estructura del programa AQUAD

3.1 Paradigmas de análisis de datos



En el análisis cualitativo de datos, la interpretación del contenido, es decir, la determinación de un significado concreto de los segmentos de contenido relevantes, es necesaria en cualquier caso en las fases críticas del proceso. Veamos primero los procesos y las funciones del programa previstas para ello en el caso del paradigma de codificación más utilizado:

El objetivo del análisis cualitativo del contenido latente, es decir, del contenido de significado de los datos no numéricos, es reducir las descripciones, explicaciones, justificaciones, notas de campo, protocolos de observación, entrevistas, etc., que suelen formularse de diversas maneras, de forma que los datos originales se resuman en afirmaciones sistemáticas sobre su significado. La forma de llevar a cabo esta tarea varía de una persona a otra y, en última instancia, depende de la pregunta de investigación de un estudio concreto. Sin embargo, existe una invariante en todos los análisis cualitativos que siguen el paradigma de la codificación: la clasificación o categorización de los datos y su representación mediante códigos. Las categorías pueden considerarse como "contenedores" en los que se colocan los datos según su significado. Se puede

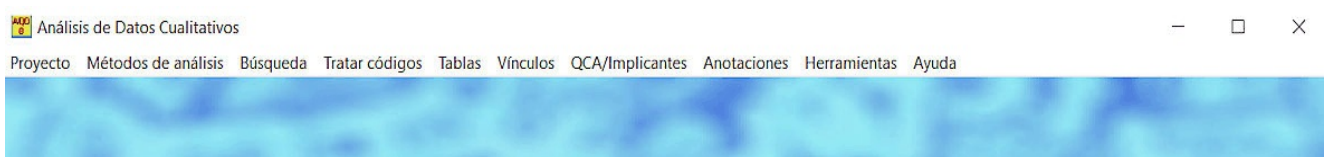
- Proceder de forma totalmente abierta e inductiva a desarrollar categorías a medida que "emergen" durante la interpretación de los datos o incluso en el curso de la recogida de datos (cf. Glaser & Strauss, 1979).
- Desarrollar categorías de forma deductiva a partir de las preguntas clave de la investigación o de las hipótesis con las que se aborda la interpretación de los datos.
- Utilizar un sistema de categorías o códigos ya disponibles en la literatura de investigación o en el propio trabajo.

AQUAD admite tanto los procesos deductivos como los inductivos, así como una combinación de ambos. Se pueden analizar todos los tipos de archivos (textos, audios, vídeos, imágenes) con los que trabaja AQUAD.

En el análisis cuantitativo del contenido manifiesto, la determinación de las frecuencias y las diferencias de frecuencia de los elementos de los datos que son críticos para la pregunta de investigación va precedida necesariamente de una fase cualitativa-interpretativa: hay que determinar de antemano lo que se va a contar después, en el caso de los textos, así las palabras clave críticas. En esta fase, las hipótesis o teorías y los constructos pertinentes a la pregunta de investigación desempeñan un papel importante en la selección. Posteriormente, los resultados cuantitativos se ordenan, es decir, se resumen en listas o cuadros temáticos. En el sentido del enfoque de "métodos mixtos", estos resultados cuantitativos pueden compararse con los resultados cualitativos de la interpretación del contenido latente. La definición cualitativa-interpretativa de palabras clave o listas de palabras relacionadas con el significado permite la evaluación automática de textos en AQUAD, pero no de archivos de audio, vídeo e imagen.

En AQUAD, la posibilidad de realizar un análisis cualitativo de acuerdo con el *paradigma de reconstrucción* paradigma como "análisis de la secuencia" según el enfoque de la hermenéutica objetiva (Oevermann, 1996; Oevermann, Allert, Konau, & Krambeck, 1979) está integrada en un módulo separado. A diferencia de los análisis de datos cualitativos según el *paradigma de codificación*, los análisis de secuencias no parten de una visión general de todo el texto y no buscan los segmentos del archivo que son significativos para la pregunta de investigación, sino que en una primera fase de generación de hipótesis se anotan todos los significados concebibles de alguna manera para cada segmento del texto (frase, parte de la frase, secuencia de palabras). En concreto, esto significa intentar contar historias en las que el segmento de texto pueda tener sentido. La multitud de interpretaciones posibles se reduce en el transcurso del análisis, se eliminan las "historias" sin sentido, es decir, contradictorias, y las "historias" con sentido y/o coherentes se reducen a una "versión" común. La generación de hipótesis es estrictamente secuencial, segmento de archivo por segmento de archivo. Sólo cuando se está convencido de que se han apuntado hipótesis significativas para la pregunta de investigación en los segmentos de datos tamizados, se utiliza el resto de los datos para confirmar o falsificar todas las hipótesis apuntadas. Sólo en esta fase se recorren los segmentos de texto restantes de forma no secuencial en busca de justificaciones para mantener o rechazar las hipótesis generadas previamente.

3.2 Módulos y funciones del análisis de datos



Las siguientes descripciones, excepto las relativas a los "Métodos de análisis", se aplican a todas las variantes de AQUAD, es decir, a los módulos de análisis de texto, audio, vídeo e imagen, ya que tras la determinación e interpretación de las unidades de significado en el material de datos, las operaciones posteriores se realizan básicamente con los códigos. Los códigos contienen la información necesaria para la localización de una unidad de significado en el archivo (comienzo y final del segmento interpretado) y el significado en sí mismo (expresado por un código): en el caso del vídeo de ejemplo "Lucy con caja" se encuentra, entre otras cosas, la codificación " 195 200Mordiendo", que expresa que el cachorro filmado muerde entre las posiciones de la imagen 195 y 200. En los demás casos, los números de línea, las horas de grabación del sonido o las coordenadas de la imagen están disponibles para la localización.

Los enfoques cualitativo, cuantitativo y objetivo-hermenéutico están disponibles como métodos de análisis en AQUAD sólo para los datos de texto. Los datos de sonido, vídeo e imagen sólo pueden dividirse en segmentos de significado de forma cualitativa-interpretativa. Por este motivo, el manual contiene secciones separadas para ellos.

3.2.1 Proyecto

En la opción "Nuevo proyecto", da un nombre a su proyecto de investigación y luego selecciona todos los archivos (textos, audios, vídeos, imágenes - dependiendo del módulo de AQUAD con el que esté trabajando) que quiera analizar en su proyecto.



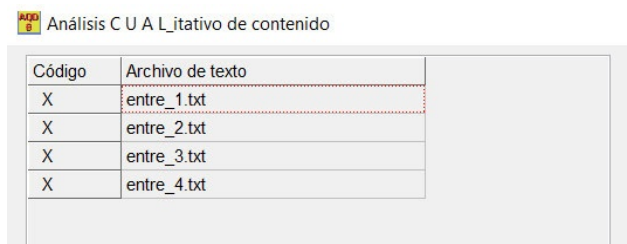
Sólo necesita la función "Abrir proyecto" si está trabajando en varios proyectos al mismo tiempo, por ejemplo, si ha combinado ciertos textos en determinadas versiones del proyecto y quiere cambiar entre los proyectos. De lo contrario, AQUAD reabre automáticamente el último proyecto abierto cada vez que lo inicia.

Con la función "END" concluye su trabajo y sale del entorno del programa AQUAD.

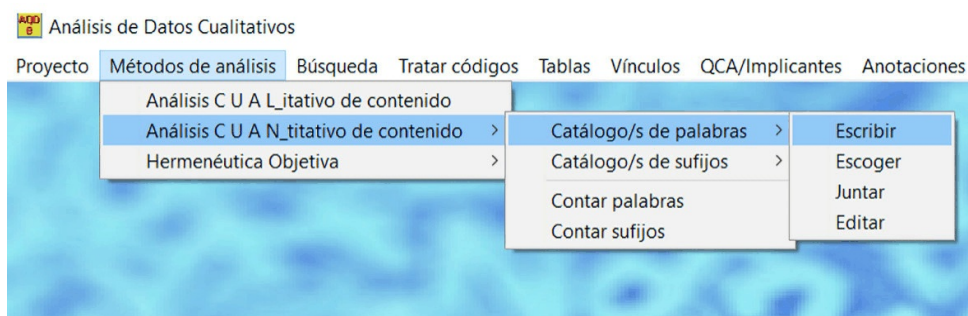
3.2.2 Métodos de análisis

Desde los inicios de una amplia recepción socio-científica de los procedimientos de análisis de texto, el análisis de contenido se ha llevado a cabo como un análisis cualitativo o cuantitativo. Ese mismo año, 1952, se publicaron artículos fundamentales de Kracauer y Berelson. Mientras que para Kracauer el análisis de contenido cualitativo intenta revelar las categorías de significado ocultas y latentes en el texto, independientemente del contenido textual manifiesto, para Berelson el análisis de contenido cuantitativo sirve para registrar sistemática, objetiva y cuantitativamente el contenido comunicativo manifiesto.

Al seleccionar "Análisis CUAL_itativo de contenido" se abre inmediatamente una pantalla para seleccionar uno de los textos del proyecto y, a continuación, la ventana de trabajo para analizar el contenido latente (véase más abajo).

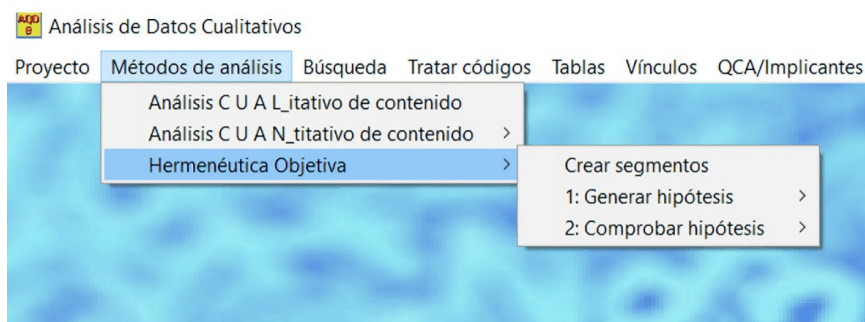


En los módulos de AQUAD para el análisis de sonido, vídeo e imagen faltan las opciones "Análisis de contenido cuantitativo" y "Hermenéutica objetiva". Aquí sólo encontrará la opción "Análisis cualitativo" y, a continuación, la opción de seleccionar un archivo de su proyecto para su posterior análisis.



La selección "Análisis CUAN_itativo de contenido" ofrece la posibilidad o, más exactamente, la exigencia de introducir los elementos críticos a contabilizar en un catálogo, de seleccionar un catálogo ya creado, de combinar catálogos o de editar un catálogo disponible. Lo mismo se aplica, mutatis mutandis, si no se quieren contar las palabras clave sino ciertos sufijos.

El análisis cuantitativo comienza entonces con el "recuento de palabras" o el "recuento de sufijos". Los resultados de estas funciones se guardan en listas y en tablas CSV ("valores separados por comas") en el subdirectorío ".\res" según diversas condiciones, por ejemplo, límites inferiores y superiores de la consideración de los elementos.

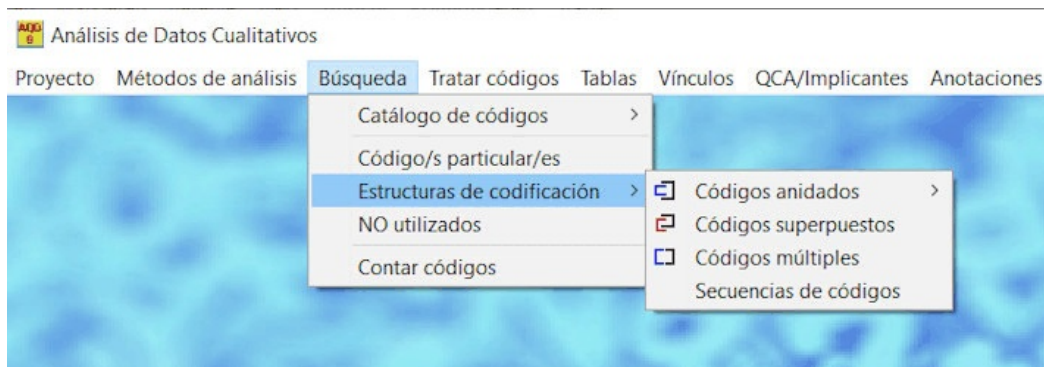


Para trabajar según las condiciones del paradigma de reconstrucción de la hermenéutica objetiva, es decir, aquí en AQUAD con el método de análisis de secuencias, primero hay que preparar las secuencias textuales. La función "Crear segmentos" ofrece desglosar el texto por adelantado en frases completas o partes de frases (cada una hasta el siguiente signo de puntuación). Durante el análisis posterior, sólo se muestra uno de estos elementos de la secuencia a la vez para generar hipótesis ("contar historias"). En una tercera variante, se puede juntar un segmento actual a la vez seleccionando cualquier número de palabras consecutivas.

En una primera fase analítica de generar hipótesis, se anotan todos los significados concebibles para cada segmento de texto. La generación de hipótesis es estrictamente secuencial, segmento de archivo por segmento de archivo. Sólo cuando se está convencido de haber anotado todas las hipótesis que son significativas para la pregunta de investigación en los segmentos de datos revisados, se utiliza el resto de los datos para confirmar o rechazar las hipótesis anotadas. En esta segunda fase de comprobación de hipótesis, se recorren los segmentos de texto restantes de forma no secuencial en busca de justificaciones para mantener o rechazar las hipótesis generadas anteriormente.

3.2.3 Búsqueda

Hay muchas razones por las que un investigador puede querer de vez en cuando revisar todos los pasajes de los archivos que tratan un mismo tema, es decir, áreas en textos, grabaciones de audio o vídeo o imágenes, que entren en la misma categoría. Una de estas razones, que desempeña un papel importante en muchos estudios descriptivos-interpretativos (véase Tesch, 1990), es detectar puntos comunes en toda la base de datos. Otra razón se encuentra en el esfuerzo por controlar la coherencia de la codificación; hay que asegurarse de que siempre se han utilizado los mismos principios.



Otras razones pueden encontrarse en la búsqueda de pruebas de dónde se ha utilizado el código correctamente y dónde se ha utilizado incorrectamente, dónde hay paralelismos entre ciertos archivos, etc. AQUAD muestra las áreas que busca en todos los archivos. Esto se hace en el orden en que aparecen en el directorio de archivos (creado cuando se crea un nuevo proyecto). Dado que se busca en un archivo a la vez, llamamos a esta búsqueda unidimensional o análisis lineal. Se diferencia de los análisis bidimensionales, llamados análisis de tablas o matrices en AQUAD, y de los análisis de enlaces complejos. Los resultados de los análisis unidimensionales pueden leerse en la pantalla, pueden imprimirse en papel o guardarse primero en el disco duro en formato txt.

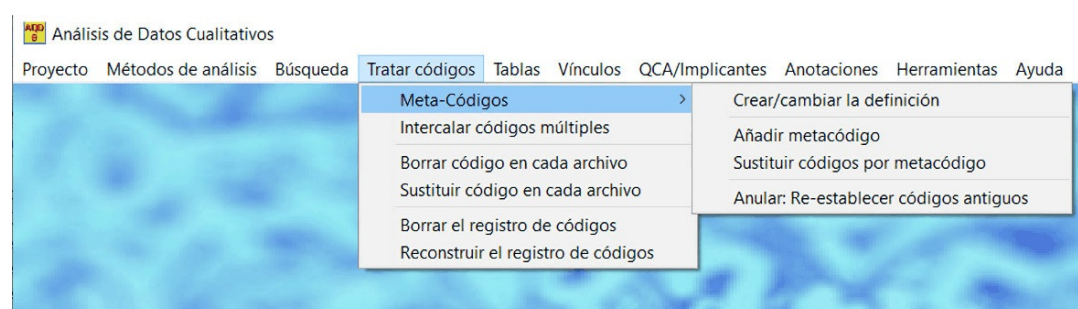
La búsqueda lineal de segmentos de archivos codificados se inicia en la opción de menú principal "Búsqueda" con la opción "código/s particular/es". Es posible buscar varios códigos (de un catálogo de códigos) al mismo tiempo. También es posible buscar determinadas estructuras de codificación, a saber, si un código aparece por encima o por debajo de otro, si su rango de significado se solapa con el de otros códigos, si una unidad de significado ha sido marcada varias veces con diferentes códigos, si un código aparece en una determinada secuencia con otros y, por último, si dichas secuencias de códigos se repiten. Además, los archivos pueden compararse para ver si determinados códigos no se han utilizado para la interpretación. Por último, se puede contar la frecuencia de los códigos y, en los archivos de vídeo y audio, la duración de determinadas escenas.

También es posible buscar varios códigos al mismo tiempo. Estos códigos se resumen en listas con la primera función ("catálogo de códigos").

En el programa "aquad_8c_graph.exe" para el análisis de material gráfico, sólo se pueden buscar códigos múltiples con la opción "estructuras de códigos", ya que en el material de la imagen la información total está disponible simultáneamente y las coordenadas de las secciones de la imagen marcadas no definen necesariamente una secuencia. En principio, podrían analizarse las inclusiones y solapamientos, pero hasta ahora no ha habido peticiones de los usuarios.

Con la información sobre la frecuencia de las codificaciones seleccionadas de un proyecto, son posibles muchos análisis cuantitativos bastante significativos, dependiendo de la pregunta de investigación. En cualquier caso, el punto de partida es contar la frecuencia de los códigos con la función "Contar códigos" del submenú "Búsqueda". El resultado se muestra en una simple lista de frecuencias de los códigos contados y puede guardarse como un archivo de texto (*.txt). Automáticamente se guarda en una tabla estructurada según códigos y archivos en el formato *_DT.csv en el subdirectorio ..res (_DT significa "tabla de datos"). Estas tablas pueden utilizarse sin problemas en el módulo "Aquad_eda.exe" para el análisis exploratorio de datos o exportarse a programas de hojas de cálculo o para otros análisis estadísticos.

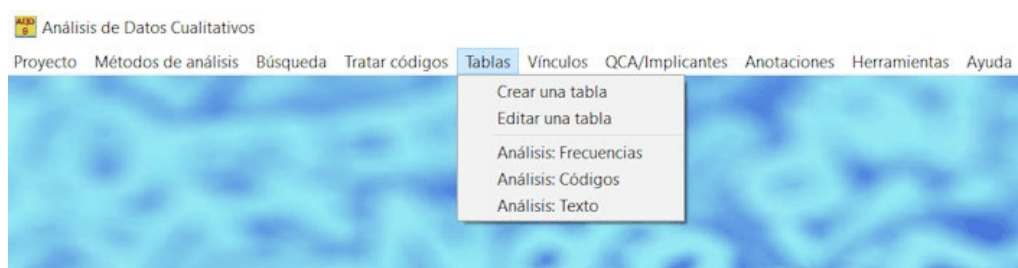
3.2.4 Tratar códigos



Se puede utilizar la función "Metacódigos" para combinar los códigos relacionados en un solo código. Para determinados códigos (por ejemplo, para un código "entrevistador"), también es posible que se añadan automáticamente códigos adicionales a posteriori (por ejemplo, el código de control "\$no contar", que excluye el segmento de texto codificado del recuento de palabras). Los códigos inadecuados pueden eliminarse de todos los archivos o sustituirse por un código más adecuado (o correctamente escrito). Cuando se codifica, AQUAD lleva un registro de los códigos que se han utilizado hasta ahora en un proyecto. Si, por ejemplo, se producen errores aquí debido a la edición directa de los archivos de código fuera de AQUAD, se puede eliminar el registro de código sin preocuparse y restaurarlo automáticamente de forma mejorada con la ayuda de los archivos de código disponibles.

3.2.5 Tablas

En este módulo se puede comparar cómo se producen determinados códigos en archivos caracterizados por diferentes "códigos de perfil", es decir, la función de tabla compara determinados segmentos de significado bajo la condición de otros códigos específicos en los archivos.



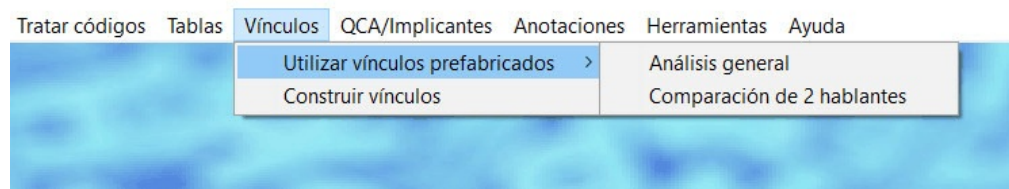
Supongamos que en el análisis de las entrevistas se registró el género de los hablantes y queremos comparar lo que los entrevistados dijeron sobre el trabajo y el ocio. Con los códigos "/femenino" y "/masculino" podríamos ahora definir las columnas, con los códigos "trabajo" y "ocio" las filas de una matriz de 2x2.

En la primera de las tres funciones de análisis posibles ("Frecuencias") para una visión rápida, sólo la frecuencia de los códigos en las celdas de la matriz. La segunda función de análisis ("codificaciones") indica las ubicaciones de los códigos en los archivos, la tercera función de análisis ("segmentos de texto") también da salida en detalle a todos los segmentos de texto que satisfacen la respectiva combinación de condiciones, es decir, las cuatro celdas de la matriz de nuestro ejemplo se llenarán con segmentos de textos en los que los hombres hablan de actividades de ocio, luego del trabajo, luego encontramos declaraciones de las mujeres primero sobre las horas de ocio, en la siguiente celda sobre el trabajo.

3.2.6 Vínculos

Este es uno de los dos módulos más importantes de AQUAD para la construcción de la teoría. Puede utilizarse para identificar vínculos sistemáticos significativos entre secciones de archivos. O, en otras palabras, AQUAD apoya la interpretación según el paradigma de codificación no sólo en la categorización de las secciones de los archivos y la organización de los datos según las categorías, sino que también facilita el descubrimiento de las relaciones entre las categorías. Se formulan expectativas de que las relaciones entre las categorías que se consideran típicas o significativas se produzcan sistemáticamente. Si uno sospecha que existen tales relaciones, debe ser capaz de poner a prueba los

datos para comprobarlas; o en palabras de Miles y Huberman (1984), Shelly y Sibert (1985) y Shelly (1986), uno debe ser capaz de poner a prueba sus hipótesis. Para eso está el módulo "Vínculos" de AQUAD. Un resultado positivo, es decir, la constatación de que determinadas combinaciones de afirmaciones se producen realmente de forma sistemática en un conjunto de datos, se consideraría una prueba de la hipótesis.



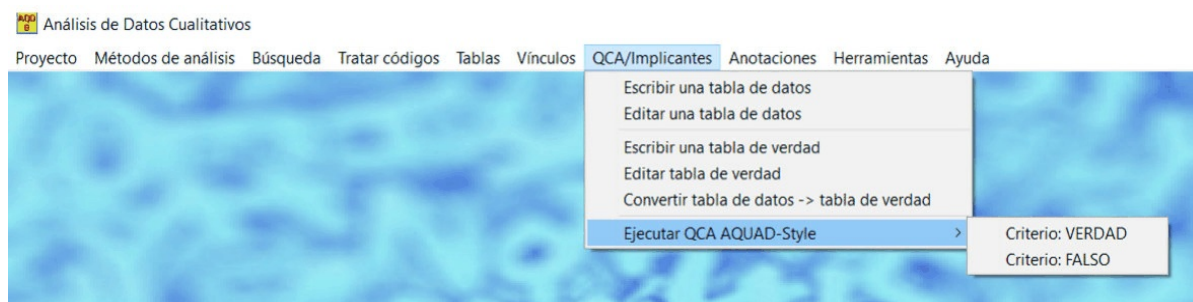
Por supuesto, primero hay que formarse una hipótesis sobre la sucesión regular de dos y más unidades de significado codificadas, por ejemplo: "Siempre que los entrevistados mencionan el tema A, pasan a hablar de B o C". A continuación, el programa comprueba si esa supuesta vinculación de códigos está presente en los datos, con qué frecuencia y dónde.

Se pueden seleccionar estructuras de enlace predefinidas y utilizar los propios códigos como "variables" o construir enlaces propios de hasta cinco códigos con los operadores lógicos "Y", "O" y "NO".

La comprobación de estos vínculos puede realizarse de forma general, sin más condiciones, o puede limitarse a la comparación de las secciones de los archivos de un hablante con las secciones de los archivos de otro hablante.

3.2.7 QCA/Implicantes

AQUAD ofrece una ayuda aún mayor con las conjeturas sobre las relaciones condicionales en los datos. Por lo general, estas hipótesis se consideran tan relacionadas con los procedimientos experimentales y estadísticos que se han evitado en gran medida en la investigación cualitativa. Sin embargo, Ragin (1987) nos recordó que el análisis de las relaciones causales tiene una venerable tradición dentro del análisis cualitativo que se remonta a John Stuart Mills (Ragin 1987, p. 36f.), y que los métodos cualitativos pueden incluso ser superiores a los análisis causales cuantitativos en algunos aspectos si no se da prioridad a la generalización sobre la complejidad (Ragin 1987, p. 54). Una discusión detallada del argumento de Ragin está fuera del alcance de este manual, pero su enfoque metodológico, el llamado "Método Booleano de Análisis Cualitativo Comparativo" o "Minimización Booleana" se explica en un capítulo separado del manual. Este método comparativo cualitativo integra -y no se limita a combinar- rasgos característicos de los diseños experimentales e interpretativos al tratar la existencia de una determinada "condición", es decir, la presencia de una determinada categoría en un conjunto de datos como una variable categórica dicotómica o "valor de verdad". La información disponible sobre el evento o la "condición" se reduce al valor "verdadero" o "falso", es decir, si el evento existe o no existe en el conjunto de datos dado. Las causas siempre se consideran combinaciones complejas de condiciones vinculadas a un "resultado" específico. Todos los datos se examinan para detectar la presencia o ausencia de todos los tipos de combinaciones posibles. Los resultados, es decir, la presencia o ausencia de una condición, se introducen en una tabla, en la que cada celda está ocupada por un cero o un uno. Mediante el uso de procedimientos algebraicos desarrollados por el matemático George Boole (1815-1864), conocidos como "lógica combinatoria" o "minimización" y el uso de "implicantes principales", se extraen conclusiones de la tabla sobre las diferentes combinaciones de condiciones de los datos en cuestión.



Resumido brevemente: Este módulo utiliza el procedimiento de minimización booleana para comparar la presencia o ausencia de una serie de códigos interrelacionados, cuando se sospecha que uno de los códigos está críticamente relacionado con ciertas configuraciones de los demás. Debe haber un número de casos lo suficientemente grande para que el resultado de esa comparación tenga sentido.

La tabla de datos que se necesita como salida puede ser escrita y editada manualmente o generada por el programa a partir de una tabla de frecuencias de códigos (véase el módulo "Editar códigos") y convertida en una tabla de valores de verdad (sí/no, verdadero/falso, 1/0). Sin embargo, también se puede crear o editar esta tabla de valores de verdad manualmente.

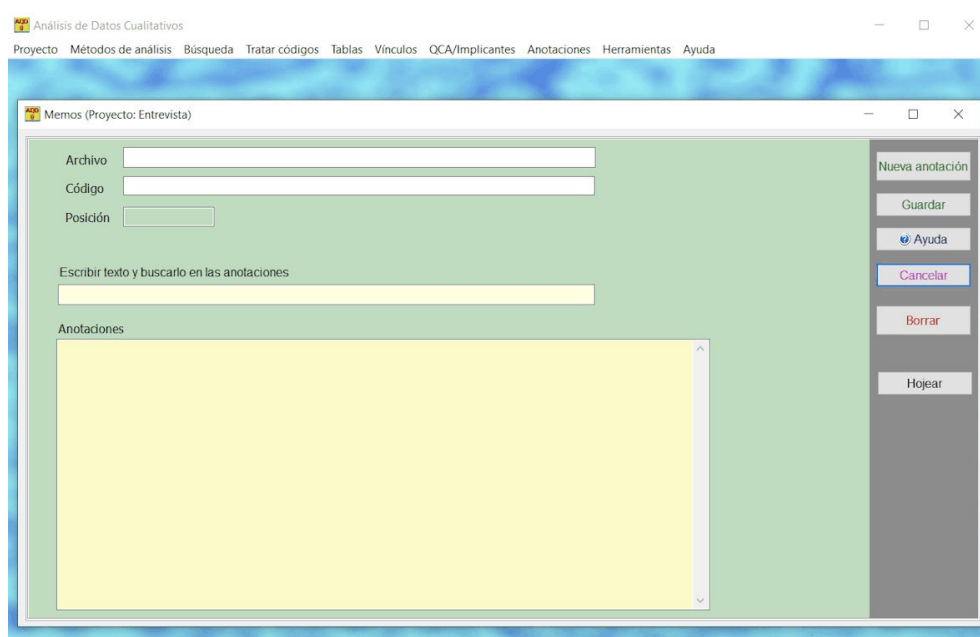
Para el análisis, se selecciona uno de los códigos hipotéticamente vinculados como criterio positivo o negativo. A continuación, el programa examina todos los casos en los que el criterio aparece con el valor "verdadero" o "falso" e indica las configuraciones de los demás códigos que se encuentran bajo esta condición.

3.2.8 Anotaciones

AQUAD ofrece dos formas posibles de acceder a sus funciones de anotaciones, dependiendo de los módulos con los que se trabaja:

1. Hay una opción "Anotaciones" en el menú principal, y
2. Hay una columna "Anotaciones" en la parte izquierda de la tabla de codificación (en las opciones de "Análisis de Contenido" del menú principal).

En el módulo "Anotaciones", siguiendo la recomendación de Glaser (1978), puede anotar inmediatamente lo que se le ocurra sobre los códigos y la codificación, qué conexiones, contradicciones, preocupaciones, excepciones, etc., le vienen a la mente mientras está ocupado interpretando un archivo y no puede seguir inmediatamente esta idea. Ante la alternativa de olvidar ideas quizá muy importantes o de perder el hilo de la interpretación de un archivo por una interrupción prolongada, tomar notas a corto plazo es un buen compromiso. AQUAD permite etiquetar estos memos con el número del archivo, la posición de un segmento relevante y el código que pueda ser importante, para que se pueda buscar específicamente la anotación más tarde. Al menos uno recuerda una u otra palabra clave en el texto de la nota, por lo que también es posible una búsqueda libre de palabras sospechosas en las notas.



3.2.9 Herramientas

El módulo de herramientas ofrece a los usuarios de la versión anterior AQUAD 7 la posibilidad de tomar archivos, codificaciones y anotaciones y procesarlos posteriormente en AQUAD 8, es decir, puede utilizar este módulo para convertir sus archivos antiguos en el formato modificado de AQUAD 8.



El autor estará encantado de asumir otras herramientas, por ejemplo, para la colaboración en el análisis de contenido en grupos o para el cálculo de elementos de imagen cuando se trabaja con el módulo de gráficos, de la versión 7 y adaptarlas a la versión 8 si se solicita.

Capítulo 4: El proceso de análisis de contenido cualitativo

En este capítulo se intenta dar una visión general de las fases típicas del análisis de datos cualitativos en unas pocas páginas. Para una mejor comprensión, se destacan los pasos individuales que caracterizan a estas fases, aunque el proceso analítico en realidad siempre se desarrolla en ciclos, es decir, uno está temporalmente involucrado, al menos "en el fondo de su mente", con actividades que caracterizan a otras fases. Además, este capítulo trata de explicar cómo se pueden utilizar las posibilidades que ofrece AQUAD. Los detalles del procedimiento práctico se explicarán con más detalle en capítulos posteriores. Este manual no puede ofrecer más que una primera visión general del procedimiento metodológico de los análisis cualitativos y de las posibilidades de utilización de los ordenadores. Para una descripción

más detallada del enfoque del análisis cualitativo, se recomienda a los lectores interesados que consulten los siguientes volúmenes:

- Se pueden encontrar contribuciones generales sobre el uso de ordenadores en la investigación cualitativa, por ejemplo, en Tesch (1990), Huber (1992), Kelle (1995), Fielding y Lee (1998) o Lissmann (2001).
- Miles y Huberman (2ª edición, 1995) ofrecen una introducción detallada al análisis interpretativo de los datos cualitativos. Se basan en numerosos ejemplos. Otros ejemplos concretos de una variedad de estudios cualitativos procedentes de la investigación educativa pueden encontrarse en Bos y Tarnai (1998) y Schratz (1993), y de la investigación psicológica en Kiegelmann (2001, 2002, 2003).
- En Mayring (2012) se puede encontrar una visión breve y clara de los principios básicos del análisis de contenido cualitativo.
- Strauss y Corbin (1990) ofrecen una introducción específica a las técnicas de análisis cualitativo de construcción de la teoría.

El siguiente resumen es una adaptación muy abreviada de una contribución propia en Huber (1992); para más detalles y otros ejemplos, consulte este volumen.



Las fases típicas del análisis de datos cualitativos son la reducción de los datos originales, la reconstrucción de las relaciones relacionadas con el contenido y la comparación de los resultados con el objetivo de la generalización. Especialmente en los estudios psicológicos, suele interesar la conclusión inductiva de los datos de una persona a las regularidades de su experiencia y comportamiento como hallazgo generalizado. También es interesante saber si se pueden encontrar puntos comunes entre varias personas o situaciones de observación, pero sólo en una fase posterior de la investigación.

- La primera fase del trabajo consiste en reducir la cantidad de datos y, al mismo tiempo, la diversidad de expresión que contienen, es decir, las formulaciones lingüísticas de los textos y, además, las formas de expresión no verbales de las grabaciones de vídeo. Para ello, se determinan segmentos de datos más o menos extensos (texto, segmentos de cintas de audio o vídeo o fragmentos de imágenes), a cada uno de los cuales se atribuye un significado definible. A continuación, se adjunta un código al segmento de datos como abreviatura o designación del significado. A continuación, estos códigos se utilizan como representantes de los segmentos de datos o de las unidades de significado de los ficheros. Básicamente, se trata de un proceso de categorización en el que las categorías exactas suelen surgir sólo durante la interpretación. Sin embargo, también pueden tomarse de un sistema de categorías ya existente. Esto depende totalmente de la orientación epistemológica del investigador.
- En la segunda fase, se reconstruyen los sistemas de significado subjetivo de los productores de datos, es decir, los entrevistados, los diaristas, los observadores, etc., a partir de estas unidades de significado. En la reconstrucción, se buscan vínculos regulares de unidades de significado dentro de los archivos que sean característicos de los productores de datos y/o de su situación.

- Por último, en la tercera fase, se desarrollan invariancias o conexiones generales comparando las percepciones de los sistemas de significado individuales obtenidas de este modo (cf. Ragin, 1987).

Es importante señalar que estas fases no son estrictamente delimitables y se suceden de forma lineal, sino que suelen solaparse y transcurrir de forma cíclica (cf. Shelly & Sibert, 1992).

Ya durante la reducción, por ejemplo, se piensa en teorías implícitas de los productores de datos, se empieza a comparar constantemente los datos; en el proceso, tal vez surja claramente un aspecto en la persona B/archivo B que se había pasado por alto en la persona A/archivo A, por lo que se vuelve a pasar por la fase de reducción al menos parcialmente allí, etc. En cada una de estas fases, siempre es importante reafirmar deductivamente la validez de las categorizaciones infiriendo las particularidades de la fase de análisis anterior y buscando las pruebas correspondientes.

4.1 La reducción de los datos cualitativos

El principio de la reducción es aparentemente sencillo, pero su aplicación resulta muy laboriosa, requiere mucho tiempo y es propensa a errores: hay que reducir el extenso material verbal a las unidades de significado que contiene.

Si los datos fuente para el análisis están disponibles en forma de grabaciones de vídeo o audio, se plantea la cuestión de si primero hay que hacer accesibles los datos para un análisis detallado transcribiéndolos, es decir, convirtiéndolos en una versión de texto. AQUAD 8 ayuda a ahorrar tiempo y dinero, ya que el software permite la reducción directa de los datos a la codificación sin los desvíos de las transcripciones. Para las secuencias de contenido crítico y su posterior presentación en informes, ensayos, etc., es aconsejable transcribirlas extracto a extracto. Esto puede hacerse con la función "Anotaciones" ("memo") de AQUAD.

La pregunta crítica en esta fase inicial del análisis cualitativo es en cambio ¿Cómo y dónde encontrar unidades de significado en los archivos? Tanto para los principiantes en el análisis cualitativo como para los que se están familiarizando con nuevas áreas de contenido, esta cuestión se plantea una y otra vez. Weber (1985) se refiere a seis formas habituales de determinar las unidades para los archivos de texto en general, a saber, elegir como unidad de análisis palabras, significados de palabras, frases, temas, secciones, el texto completo (posible, por ejemplo, para titulares, resúmenes, cartas cortas al editor, etc.). Sin embargo, estas determinaciones no pueden hacerse de forma mecánica, sino que requieren decisiones cualitativas previas. En el caso de los vídeos, son necesarias decisiones adicionales para la información no lingüística, no verbal, pero también a nivel de palabras individuales, pero sobre todo en el caso de las alternativas más complejas, como los significados de las palabras, las frases, etc., como unidades de análisis, debe haber una constatación o hipótesis previa de que la unidad seleccionada es relevante en el sentido de la pregunta de investigación. Además, es necesario tomar decisiones cualitativas adicionales, por ejemplo, sobre qué palabras se utilizan como sinónimos o qué expresiones idiomáticas tienen también el mismo significado para los productores de los textos que se van a analizar.

Con la determinación de la unidad de significado se establece una diferencia esencial entre los enfoques cuantitativos y cualitativos desde los inicios de una amplia recepción socio-científica de los procedimientos de análisis de textos. Ese mismo año, 1952, se publicaron artículos fundamentales de Berelson y Kracauer. Mientras que para Berelson el análisis de contenido sirve para registrar de forma sistemática, objetiva y cuantitativa el contenido manifiesto de la comunicación, el análisis de contenido cualitativo, según Kracauer, intenta revelar las categorías de significado ocultas y latentes en el texto, independientemente del contenido textual manifiesto. Ambos autores adoptan posiciones opuestas en un debate que Thomas y Znaniecki (1918) habían desencadenado 35 años antes con un análisis

sociológico ya clásico de las cartas de los emigrantes polacos de América. La cuestión de la unidad analítica adecuada -palabras frente a significados- se confunde obviamente con el objetivo del análisis textual.

En el plano procesal, los dos enfoques no son fundamentalmente excluyentes. Especialmente con apoyo informático, se pueden utilizar provechosamente diferentes procedimientos del enfoque cuantitativo para el análisis cualitativo de textos e imágenes como apoyo a los propios esfuerzos de interpretación.

Básicamente, a la hora de seleccionar unidades de significado para el análisis de datos cualitativos, siempre hay que tener en cuenta que, independientemente de si se quiere comprender las experiencias y acciones de las personas desde su punto de vista, es decir, desde sus propias representaciones verbales, o si se quiere explicar algunas de sus acciones registradas remontándolas al marco de referencia de las teorías subjetivas de estas personas, hay que inferir las unidades analíticas durante la interpretación de los datos. Sólo este enfoque inductivo garantiza que se acceda a la visión subjetiva del mundo de los interlocutores o de los observados y no sólo a algunos aspectos parciales de la misma, posiblemente desvinculados del contexto subjetivo de significado, queden atascados en la parrilla analítica del investigador. Esta estrategia corresponde al enfoque de Glaser y Strauss (1979) de la "teoría fundamentada".

Sin embargo, este enfoque requiere un trabajo considerable en cuanto se tiene más de un conjunto de datos para analizar. Dado que se desea que los resultados individuales sean comparables en una fase más avanzada del análisis, hay que examinar -y normalmente varias veces- las unidades de significado y su codificación en todos los archivos para comprobar su coherencia y posiblemente modificarlos. Por ello, Miles y Huberman (1994) han propuesto un compromiso en el que se establece un marco de orientación muy general, no específico del contenido, para la búsqueda de unidades de significado antes de leer los textos/ver los archivos, dentro del cual se deciden luego las unidades específicas según las condiciones del archivo individual. Este compromiso parece justificado en la medida en que la reducción de datos también comienza al principio de una investigación, de hecho ya en su planificación.

4.2 Esquema del proceso de análisis de datos

(1) Es indispensable familiarizarse primero con el tema investigado desde la perspectiva del entrevistado o del observado, especialmente si uno mismo no estuvo presente en la situación de la entrevista, por ejemplo, si habló con el entrevistado como entrevistador, o si no participó en las grabaciones de vídeo. En otras palabras, no hay que intentar desarrollar un sistema diferenciado de categorías con el primer conjunto de datos. Sin duda, sólo sería válida para este conjunto de datos actual y tendría que revisarse fundamentalmente al analizar el segundo conjunto de datos. En su lugar, es aconsejable seleccionar aleatoriamente entre 8 y 10 transcripciones o grabaciones de un número mayor de conjuntos de datos o, de forma análoga a la selección de casos en los análisis de casos individuales, escoger selectivamente algunos conjuntos de datos de todo el material de datos según el principio de "muestreo teórico".

Por ejemplo, si queremos analizar las grabaciones de la enseñanza de las matemáticas en distintos centros escolares, cabe esperar que haya diferencias en función del grado, el tipo de centro y la experiencia docente de los profesores observados. Según estos criterios, seleccionaríamos unos cuantos conjuntos de datos y elaboraríamos una primera orientación a partir de ellos. Sin embargo, tal vez hayamos observado desde el principio clases en las que la enseñanza y el aprendizaje se realizan según métodos diferentes. Entonces (también) utilizaríamos conjuntos de datos según esta característica para la orientación.

Esta primera fase modifica, obviamente, la recomendación de Miles y Huberman (1994; véase más arriba) de establecerse en un marco de orientación general, no específico del contenido, para buscar unidades de significado incluso antes de leer los textos o revisar los archivos. Recomendamos aquí que este marco de orientación, que puede

estar disponible de antemano en forma de esquema debido a la pregunta de investigación, se compruebe y diferencie con la ayuda de una muestra de datos.

(2) Si seleccionamos los archivos según el principio de "muestreo teórico", ya tenemos que pensar en detalle sobre las características generales de los conjuntos de datos. Consideramos lo que determina el perfil de los conjuntos de datos individuales. En el caso de los datos procedentes de las "observaciones de clases", por ejemplo, por la edad de los alumnos, el tipo de escuela, la experiencia profesional de los profesores, quizás también por el tamaño de la clase, la orientación didáctica de los profesores, etc. Estas características deben anotarse en una nota. Deberíamos registrar estas características en una nota y utilizarlas posteriormente como "códigos de perfil" (o "singulares", es decir, categorías utilizadas una sola vez por conjunto de datos para su caracterización global) para la codificación. Lo mejor es insertar estos códigos de perfil justo al principio del archivo.

Si, por el contrario, nos limitamos a elegir unos cuantos conjuntos de datos al azar para orientarnos, sería conveniente considerar ahora de antemano qué características de las personas entrevistadas/observadas y de la situación respectiva en relación con la pregunta de investigación del estudio podrían llegar a ser significativas como "características del perfil" para el análisis posterior.

(3) En primer lugar, leemos/vemos la selección de datos (véase el punto 1) sin una codificación elaborada e intentamos comprender el contexto general, distinguir las secciones esenciales, encontrar puntos comunes y contrastes y, por tanto, definir nuestro marco de orientación. Es fundamental que dejemos constancia de nuestras interpretaciones e ideas en memos.

Al final de este paso es aconsejable escribir un breve (!) resumen temático, preferiblemente de nuevo en forma de nota. En él anotamos nuestra "primera impresión" de lo que supone este conjunto de datos (entrevista, grabación, etc.), lo que consideramos que es el significado central de este conjunto de datos en esta primera fase del análisis. También deberíamos esbozar el marco de orientación -siempre preliminar- en una nota aparte.

(4) Entonces tenemos que tomar una decisión importante sobre nuestra estrategia interpretativa básica: ¿Queremos, como ya se ha sugerido en el punto 3 del resumen temático, buscar las unidades de significado más completas posibles, codificarlas y luego diferenciarlas paso a paso y en pases repetidos? ¿O queremos prestar atención a todos los detalles desde el principio, limitados a lo sumo por el marco de orientación general, marcarlos con la ayuda de códigos y luego resumir los numerosos detalles en categorías generales en pases de codificación posteriores para que los conjuntos de datos sean comparables? La primera estrategia va de lo general a lo particular ("diferenciación"), la segunda estrategia va de lo particular a lo general en los conjuntos de datos ("generalización").

Por supuesto, también depende de la pregunta de investigación y de la experiencia de los investigadores con el tema qué estrategia es preferible. En principio, sin embargo, desaconsejamos a los principiantes en el análisis de datos cualitativos la estrategia de generalización, en la que primero se examinan todos los rasgos específicos y luego, en pasos posteriores, se combinan en categorías más abstractas y generales. De este modo, uno suele mantenerse con éxito alejado de consideraciones analíticas más profundas, inicialmente difíciles, orientadas a la teoría, y corre el riesgo de entregarse a una ocupación a menudo superficial durante mucho tiempo. Se producen códigos, se está activo y, por tanto, se está tranquilo, aunque luego haya que darse cuenta de que las actividades no eran necesariamente intencionadas. En un seminario, por ejemplo, un grupo de estudiantes asignó códigos para las palabras clave (según los estudiantes) del texto, básicamente duplicando los datos. En otro caso, un principiante generó unos 1.500 (¡no es una errata!) códigos para la interpretación de sus entrevistas y luego, por supuesto, dejó de ver el bosque por los árboles.

Por lo tanto, nos parece más prometedor, especialmente para los principiantes, tomarse muy en serio desde el principio el principio de "comparación constante" en el enfoque de la "teoría fundamentada" (Glaser y Strauss, 1979) y buscar las similitudes y diferencias primero dentro del individuo, y luego entre los conjuntos de datos. Con esta orientación, uno se asegura en gran medida de no perderse en los detalles superficiales, sino de trabajar decididamente hacia la comprensión de las estructuras de significado. La siguiente tabla compara de nuevo los enfoques:

Estrategia:	Diferenciación	Generalización
Codificación comienza con	Buscar categorías generales	Buscar aspectos específicos
Procedimiento adicional	Diferenciación para descubrir diferencias específicas	Generalización para identificar aspectos comunes

Sin embargo, ambos enfoques no se excluyen mutuamente, sino que son dependientes entre sí. La interpretación de los datos cualitativos no es un proceso lineal sino cíclico. Shelly y Sibert (1992) ya han descrito este ciclo, basándose en las observaciones de Dewey sobre la conexión entre el razonamiento inductivo y el deductivo:

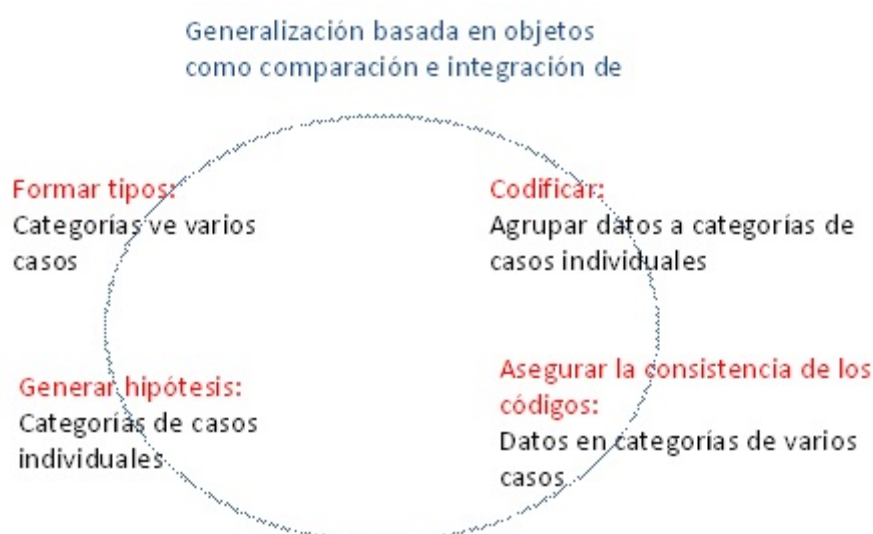


Sin embargo, desde esta perspectiva, la recomendación para los principiantes de empezar por las categorías generales debe contrarrestarse con el hecho de que uno tiene más posibilidades de descubrir algo nuevo en los datos si interpreta pequeños segmentos de datos, línea por línea, por así decirlo, pero a riesgo de pasar por alto lo general y lo esencial.

Si se sigue nuestra recomendación de abordar los datos con una estrategia de diferenciación gradual, el acento se pone inicialmente en la "base de conocimiento interpretativo". Este conocimiento interpretativo puede, por supuesto, haberse desarrollado de forma inductiva a partir de la experiencia sobre el terreno o de una primera revisión del material de datos, por ejemplo en forma de algunas categorías muy generales (estudio en un jardín de infancia: "La profesora dirige la atención de los niños"; "Les apoya en las dificultades"; "Fomenta la verbalización", etc.) o en forma de hipótesis de trabajo que ya fueron decisivas en la formulación de la pregunta de investigación. Necesariamente, la base de datos se diferencia cuando analizamos grabaciones de varias personas y quizás de diferentes situaciones.

Si aplicamos la estrategia de la generalización, es decir, si comenzamos el análisis con los significados específicos que aparecen en los datos desde una perspectiva interpretativa, las señales de las interpretaciones específicas al conocimiento interpretativo general se vuelven particularmente valiosas.

La regla básica de la "teoría fundamentada", la "comparación constante", fue caracterizada como una guía para la generalización. Estos procesos de comparación también son cíclicos, como han señalado Shelly y Sibert (1992). Siguiendo su clasificación, la comparación constante y sus objetivos pueden describirse como sigue en el curso del análisis:



(5) Tras la decisión básica sobre cómo proceder con el análisis, se concretan las consideraciones: se empieza a codificar una muestra de textos. Al hacerlo, necesariamente se pasará una y otra vez de comparar e integrar datos individuales en el conjunto de datos único a comparar varios conjuntos de datos en las categorías generadas en el proceso. De este modo, se avanza en ciclos de análisis más pequeños y menos amplios dentro del proceso cíclico, inicialmente para garantizar la coherencia de la codificación.

(6) Las reglas de codificación así obtenidas se aplican finalmente al total de los archivos disponibles. En el proceso, uno se encontrará regularmente con contradicciones y excepciones que provocan la modificación de las reglas de codificación. En algunos casos, estas percepciones pueden hacer necesario abandonar el ciclo de reducción/interpretación de datos y volver a entrar en la fase de recogida de datos en un amplio movimiento cíclico.

(7) A lo largo del proceso irán surgiendo memos de apoyo a la generación y tipificación de hipótesis. A partir de estas incursiones y de las percepciones obtenidas en el curso de la codificación, el análisis en una fase avanzada se centra principalmente en dilucidar las relaciones sistemáticas (de significado) entre las categorías críticas. Esto se describió en el punto 4 como la integración de categorías en "hipótesis"; por ejemplo, "Cuando el educador E observa el evento A o B, aplica la opción de acción X". O bien: "Cuando el profesor P habla de la falta de motivación de sus alumnos, menciona la influencia de los medios de comunicación y el papel de los padres".

Las experiencias en esta fase, por ejemplo de contradicciones en las interpretaciones, también pueden ser la ocasión para un bucle analítico de vuelta a través de las etapas de codificación y garantía de coherencia (véase el

punto 4). También puede ser necesario repetir la recogida de datos con preguntas modificadas, situaciones de observación, etc.

(8) Asumiendo un número suficiente de casos, se puede finalmente intentar resumir los casos individuales en tipos sobre la base de la comparación e integración (teóricamente justificada) de las categorías seleccionadas. Como resultado, obtenemos una diferenciación de los casos según las combinaciones típicas de características.

(9) No debe faltar aquí una nota general que se aplica al proceso de análisis cualitativo en su conjunto: el análisis cualitativo tiene que ver con la "calidad" de los acontecimientos y estados captados en los datos. Así, por regla general, las codificaciones no sólo deben captar de forma neutral el hecho de que se produzcan determinados acontecimientos, sino que también deben representar la calidad de los mismos. Por ejemplo, en un estudio de entrevistas biográficas con adolescentes, no sería suficiente para la generación de hipótesis y la formación de tipos el mero hecho de marcar con un código "padres" los segmentos del archivo en los que se habla de sus padres. ¿Qué jóvenes no hablan de sus padres?

Una característica que aparece igual en todas ellas sin variación no puede contribuir a la "comparación constante" y, por lo tanto, no es adecuada para otras distinciones y explicaciones. Así que tenemos que matizar la charla sobre los padres desde el principio: ¿Cómo hablan los jóvenes de sus padres? Unas pocas gradaciones diferenciadas, por ejemplo, positivo / indiferente / negativo, suelen ser suficientes al principio del análisis para buscar los puntos comunes y las diferencias.

Es un rasgo humanamente simpático cuando los principiantes del análisis cualitativo se retraen con tales evaluaciones, especialmente de las características o declaraciones cercanas a la persona, cuando evitan las atribuciones de valor precipitadas ("positivo"; "negativo") - pero el análisis cualitativo trata precisamente de las calificaciones. Así pues, a más tardar cuando se compruebe la coherencia de las codificaciones y se establezcan las reglas de codificación, deberán examinarse los códigos "neutros", no calificados, para ver si deben ser motivo de pasar por un bucle hacia atrás en el proceso de interpretación y comparar dentro de los archivos y entre ellos para ver si se ponen de manifiesto las diferencias de calidad de los segmentos de datos marcados con ellos.

4.3 Elaboración de códigos para unidades de significado

Los siguientes consejos deben entenderse como heurísticos, es decir, como pistas generales sobre cómo identificar unidades de significado y encontrar códigos adecuados, pero no como reglas que conduzcan paso a paso a la meta. Hay tres estrategias posibles. Aquad ofrece una valiosa ayuda, especialmente cuando se utilizan las dos primeras posibilidades y sus variantes que se presentan a continuación.

1. Cuando se buscan categorías, se recorre un archivo para ver si los acontecimientos o las declaraciones sobre acontecimientos, situaciones, personas, así como las opiniones expresadas, las ideas, etc. que contiene pueden atribuirse a un concepto supraordenado.
2. Al buscar secuencias, se presta atención a la representación de las conexiones en los eventos o enunciados, es decir, se buscan unidades de significado mucho más amplias que en una búsqueda categórica.
3. Cuando se buscan temas, se requiere la mayor abstracción; en algunos casos, todo un conjunto de datos se reduce a qué tema se aborda en él.

Cresswell (2012) proporciona ejemplos muy instructivos de estudios cualitativos con diferentes antecedentes teóricos, a saber, análisis según los enfoques narrativo, fenomenológico y etnográfico, así como los enfoques de la

teoría fundamentada y el estudio de casos. Saldaña (2016) ofrece una visión general y consejos concretos para desarrollar códigos específicos de contenido en los análisis según el paradigma de codificación.

4.3.1 Desarrollo de códigos categóricos

¿Las categorías o los códigos correspondientes deben describir, interpretar o explicar (véase Miles y Huberman, 1994, p. 56)? Por supuesto, no hay que esperar la ayuda de un ordenador en esta decisión. Metodológicamente, el investigador debe decidir sobre su propia responsabilidad. El ordenador sólo ayuda a documentar todas las decisiones, a ordenarlas, a revisarlas si es necesario sin mucho esfuerzo. Hay tres posibles soluciones al primer problema:

Uso de sistemas de categorías predefinidas

En el caso de que un estudio cuantitativo no requiera la construcción de una teoría basada en el objeto, es decir, que no se quiera desarrollar sistemas explicativos subjetivos a partir de los datos disponibles, el principio se describe rápidamente: De este modo, se puede recurrir a los sistemas de categorías ya disponibles y reducir los datos propios según los esquemas de interpretación contenidos en estos sistemas. Estos sistemas de categorías están "disponibles", por ejemplo, a partir de los propios estudios anteriores; también pueden encontrarse en la literatura empírica o en los análisis teóricos del área de contenido de interés.

Si se utilizan sistemas de categorías dados, hay que decidir qué secciones de los datos disponibles corresponden a qué ejemplos de definición de las categorías del sistema y, a continuación, registrar la ubicación y la codificación con la ayuda del ordenador. Dado que algunos enunciados no se pueden asignar de forma inequívoca y que a veces se producen incoherencias en el proceso de interpretación, una función de búsqueda de secciones del archivo codificadas específicamente ("Buscar"; "Códigos específicos") ofrece una gran ayuda para controlar la reducción de datos. Tras introducir un código crítico, esta función devuelve rápidamente y sin esfuerzo todas las secciones de archivo asignadas en todos los conjuntos de datos analizados hasta ese momento.

Sin embargo, con esta forma de control de la fiabilidad de las codificaciones, inicialmente sólo se encuentran los segmentos de datos asignados erróneamente a una categoría, pero no los pasajes asignados erróneamente no a ésta, sino a otra categoría. Para detectar este error, se podría recorrer inmediatamente al menos el conjunto de segmentos de datos de categorías similares "susceptibles" de asignaciones erróneas. Además, se puede - sólo cuando se analizan los textos (!) - hacer uso de las posibilidades de complementación mutua de la definición de las unidades de significado manifiestas y latentes. Para ello se utilizan tres estrategias de diferente complejidad:

1. Se definen las palabras clave como indicadores manifiestos de un contenido de significado crítico y se busca su aparición en todo el texto. A continuación, se puede evaluar en cada caso si el pasaje de texto en el que se encuentra una palabra clave no debe asignarse a la categoría de interés.
2. Las palabras clave se resumen en un léxico de análisis o catálogo de palabras, es decir, se introduce una lista de palabras clave para buscar pasajes críticos del texto. De este modo, se obtiene la información que se busca en una sola pasada por los textos.
3. En AQUAD, se ofrecen las dos primeras posibilidades, pero se vinculan automáticamente a la tercera variante, la función Palabras clave en contexto (Popko 1980). Esto permite buscar automáticamente una o varias palabras clave seguidas en todos los textos de un estudio. El resultado es una impresión de todas las líneas en las que aparece la palabra buscada.

Categorización basada en hipótesis

A la hora de determinar y especificar el contenido de las unidades de significado en los archivos que se van a analizar, uno se guía por hipótesis sobre el área de contenido. A partir de las hipótesis, se intenta determinar las categorías y las reglas de codificación de los datos.

Por ejemplo, el enfoque de Marcelo (véase p. 17) para clasificar las entrevistas con los profesores noveles se basó en una hipótesis. Un modelo teórico de socialización profesional sirvió de marco para diseñar un sistema preliminar de categorías. La creciente familiaridad con los textos de las entrevistas y el mayor conocimiento de la visión subjetiva del mundo de los profesores noveles hizo que se descartaran algunas de estas categorías por considerarlas poco adecuadas, mientras que otras se descubrieron por primera vez y se incluyeron de nuevo en el sistema. A medida que avanzaba el análisis, algunas de estas categorías se vincularon a unidades mayores en el sentido de las teorías subjetivas de los participantes (cf. Huber y Marcelo, 1992), que luego, a su vez, tuvieron que ser comprobadas en todo el material.

Lo mejor es definir categorías concretas para cada una de las hipótesis cuando se planifica la recogida de datos, por ejemplo al formular las preguntas guía para una entrevista, así como antes del análisis de los datos (parte deductiva). A continuación, se procede con estas categorías como se ha descrito anteriormente. Por otra parte, durante el análisis de los datos es seguro que se descubren segmentos en los archivos que no pueden asignarse a las categorías predefinidas. A continuación, se marcan estas unidades de significado y se desarrollan inductivamente categorías que las vinculan a las hipótesis específicas del tema. Por supuesto, en este procedimiento no se pueden descartar las modificaciones del marco de orientación original. Los grados de libertad y las exigencias de la propia actuación interpretativa aumentan con esta estrategia, así como las posibilidades de hacer justicia a las perspectivas específicas de los sujetos de la investigación en el análisis.

Incluso con esta estrategia, no se puede prescindir del "método de comparación constante" que es característico del enfoque de la construcción de la teoría basada en el objeto (Glaser & Strauss, 1967, 1979; Strauss & Corbin, 1990). Como se ha descrito anteriormente, el proceso central en la "comparación constante" es poner a prueba cada inferencia inductiva a partir de los detalles de la base de datos a principios más generales, aquí ahora categorías de reducción de datos, mediante inferencias deductivas sobre la base de datos.

A medida que aumentan las exigencias de la capacidad interpretativa del investigador, también aumenta la importancia del apoyo informático en este procedimiento. Cuando se comprueban las correlaciones entre categorías, se tropieza con los límites de las funciones de búsqueda simples. A partir de aquí, los análisis cualitativos requieren Aquad, que permite el razonamiento deductivo sobre la base de la programación lógica (Tesch 1990; Shelly & Sibert 1992).

1. Las posibilidades de búsqueda de segmentos de archivos codificados y -en el caso de los textos como base de datos- de palabras clave (incluidos los léxicos de búsqueda y la función KWIC) deben utilizarse aquí para la comprobación de la coherencia de las codificaciones dentro de los archivos individuales y en todos los archivos, al igual que cuando se utilizan sistemas de categorías determinados (véase más arriba).
2. A partir de las hipótesis que establecen el marco para el desarrollo de las categorías, surgen indicaciones de ciertas relaciones de las categorías. Aquad proporciona funciones para la reducción de datos basada en hipótesis para comprobar dichas relaciones. Esto puede utilizarse, por ejemplo, para determinar la relación de superordinación/subordinación de categorías individuales, secuencias de ciertas categorías o grupos de ciertas categorías. El factor común aquí es que no sólo hay que vigilar una categoría o los segmentos de

archivo representados por ella, sino dos o más categorías y la relación definida entre ellas - ¡que también puede incluir negaciones! También debe registrarse, por supuesto, el incumplimiento de la presunción de que existen determinadas relaciones entre unidades de significado.

Categorización para construir una teoría

La forma más exigente de reducción de datos cualitativos prescinde de todas las reglas de reducción predefinidas en forma de sistemas de categorías y de la estructuración del proceso de análisis mediante conceptos marco hipotéticos. Además, el investigador debe prestar atención a sus propias experiencias subjetivas, presuposiciones o prejuicios con respecto a la experiencia y el comportamiento de los sujetos de la investigación y tratar de evitar las solidificaciones prematuras comparando constantemente los principios de reducción que surgen en el proceso de análisis con los hechos o declaraciones de los archivos. La sensibilidad a las lecturas que uno mismo hace de los datos es especialmente necesaria porque los procedimientos de análisis de contenido se aplican generalmente a datos que suelen evaluarse a gran distancia del momento y el contexto de su creación (Fischer, 1982).

La información sobre los productores de los datos (personas que actúan, hablantes o escritores) y su contexto debe obtenerse de datos adicionales o del propio archivo. Este enfoque promete acercar al investigador al objetivo de ver el mundo de los sujetos de la investigación a través de sus propios ojos y comprenderlo desde su propia perspectiva.

Si, en este proceso, se intenta no sólo describir las visiones subjetivas del mundo, sino también ordenarlas con la ayuda de términos adecuados y reconstruir conexiones sistemáticas, entonces se lucha por lo que Glaser y Strauss (1967; 1979) han llamado el descubrimiento de una "teoría fundamentada" o "teoría anclada". La palabra "descubrimiento" contiene la diferencia esencial con los métodos que buscan confirmar teorías dadas. Strauss y Corbin (1990) señalan que las teorías basadas en el objeto se desarrollan en el compromiso analítico con el objeto, por lo que no se parte de la teoría para demostrarla a partir de los datos específicos del objeto, sino de los fenómenos observados y registrados que se quieren comprender y explicar mejor.

Este enfoque requiere las mayores capacidades interpretativas, por lo que presenta las mayores dificultades para los principiantes. Los sistemas de categorías contienen reglas interpretativas claras, implícitas en la elección de ejemplos para las categorías o explícitas en las definiciones de las mismas; la categorización guiada por hipótesis puede, al menos, basarse en una orientación general sobre lo que hay que buscar en los datos. La categorización para la construcción de la teoría, en cambio, debe descubrir primero de qué tratan los datos, de qué hablan los textos. Los principiantes tienden a arriesgarse lo menos posible a cometer errores y a interpretar lo menos posible (véase más arriba). En casos extremos, esta tendencia lleva a leer los datos de texto (los más utilizados) en busca de palabras relevantes, lo que es comparable al uso de palabras clave al categorizar con sistemas de categorías ya creados. Los pasajes del texto que contienen la palabra crítica se marcan como la ubicación de una "unidad de significado" y se etiquetan con un código. En este procedimiento, sin embargo, el código apenas puede representar el significado subjetivo, sino el hecho de que una determinada palabra aparece en el pasaje del texto designado.

Esta tendencia es especialmente pronunciada cuando no se conocen todavía las posibilidades de interpretación tentativa y de revisión sin esfuerzo del análisis de contenido asistido por ordenador. Con AQUAD, uno no es castigado por "jugar" con sus propias ideas mediante un trabajo de revisión poco razonable si luego encuentra contradicciones; al contrario, se le anima a interpretar, ya que los códigos son necesarios desde el principio para designar los segmentos del texto y porque los cambios, los resúmenes, las diferenciaciones pueden hacerse sin mucho esfuerzo (cf. Tesch 1992).

Por lo tanto, como regla general para la categorización que descubre la teoría, se puede adoptar lo siguiente: buscar en los datos unidades de significado que sean, por un lado, tan grandes como sea posible para que haya algo que interpretar, y, por otro lado, tan pequeñas como sea necesario para que los contenidos incoherentes no estén representados por el mismo código. Esto describe la estrategia de diferenciación recomendada anteriormente. AQUAD apoya esta estrategia porque no impone ninguna restricción a la definición de las unidades de significado; en concreto, varias unidades pueden solaparse y, a la inversa, se puede codificar la misma unidad varias veces.

Si uno está familiarizado con el área temática abordada en los datos y con el método, también puede dar la vuelta a esta regla, es decir, partir de las unidades de significado más pequeñas y seguir una estrategia de generalización. Con estas condiciones previas, en primer lugar, uno no se atasca tan fácilmente en los detalles y no se pierde la esencia de los datos. En segundo lugar, se sabe que la codificación detallada puede vincularse automáticamente a unidades mayores con la ayuda de funciones de metacodificación asistida por ordenador. Comparando constantemente los segmentos de archivo codificados, se determinan inductivamente las dimensiones de las características que circunscriben y, por último, las categorías generales a las que deben asignarse. AQUAD también es compatible con esta estrategia, ya que se puede buscar muy rápidamente, a partir de una codificación específica, segmentos de datos contrarios, similares, dependientes o definidos de otro modo. A partir de esta(s) relación(es), se puede intentar subordinar las codificaciones parciales a una categoría superior.

4.3.2 Desarrollo de códigos secuenciales

En la reducción categorial, tratamos de distinguir los segmentos de los archivos de tal manera que se les pueda atribuir significados claramente distinguibles y mutuamente excluyentes. Sin embargo, la asignación clara de significados a las categorías no implica que los segmentos de los archivos deban ser también distintos entre sí. En función de los hábitos expresivos de los productores de datos y de las lecturas de los investigadores, los segmentos de archivos asignados a las distintas categorías suelen solaparse. En casos extremos, un mismo segmento puede ser asignado a varias categorías.

La estrategia de reducción secuencial comienza con estos vínculos en los datos. Busca una conexión específica de significados, marca la ocurrencia de tales enlaces y designa los pasajes de archivos correspondientes con un código - como con la reducción categórica. Sin embargo, el código representa aquí una secuencia definida de significados. En este caso, la unidad de significado del archivo es el pasaje completo en el que se describe esta conexión.

¿Qué estrategias hay ahora, más allá de la diferenciación de categorías y subcategorías, para descubrir conexiones de significado en los datos? A continuación, distinguimos entre las estrategias de búsqueda de secuencias simples y la definición y búsqueda de patrones de secuencias complejas.

Búsqueda de secuencias simples

Además de la búsqueda de secuencias jerárquicas de categorías superordinadas y subordinadas, que Strauss y Corbin (1990) recomiendan para la reducción secuencial, es posible toda una serie de otras secuencias simples con vistas a la sistemática gramaticolingüística. La utilidad de cada una de ellas para la reducción de datos depende, por supuesto, de la pregunta de investigación. Los textos suelen buscar secuencias causales (también con palabras clave como "debido a", "porque", "para", etc.), secuencias temporales ("durante", "entonces", "antes", etc.), secuencias concesivas (restricciones positivas como "no ... sin embargo" o restricciones negativas del tipo "efectivamente... pero"), secuencias condicionales ("si... entonces"), secuencias finales ("para que", "con el fin de", "para que", etc.), secuencias comparativas, secuencias modales y secuencias de precisión.

Ni la enumeración de estos tipos de secuencias ni las posibles palabras clave enumeradas aquí pretenden ser completas. Sólo pretenden estimular los propios intentos de reducir los datos según las secuencias de significado de acuerdo con la pregunta de investigación. En este contexto, hay que destacar la función del ordenador o del software como herramienta útil para la interpretación, pero no como agente de análisis cualitativo. Por muy útiles que resulten las funciones de búsqueda de palabras en este enfoque, las posibles aplicaciones son limitadas. A menos que se analicen documentos cuidadosamente formulados, el rendimiento de los léxicos de búsqueda de secuencias específicas de significado suele ser limitado. En una entrevista libre, suele ser difícil determinar las estructuras de las frases. ¿Dónde está la cláusula principal, dónde la subordinada? En consecuencia, las conjunciones críticas para la definición de la secuencia pueden ser pensadas, pero no pronunciadas. Además, el uso de conjunciones y frases críticas es gramaticalmente incorrecto en muchos hablantes. Así, la búsqueda de palabras clave o de léxicos completos de dichas palabras no sustituye a la interpretación empática, pero puede proporcionar una heurística útil en este trabajo.

Búsqueda de patrones de secuencias complejas

A la hora de interpretar los datos, uno puede inspirarse en las características estructurales o procesales del objeto de estudio y buscar secuencias de significado más complejas que las señaladas anteriormente. Por ejemplo, para reconstruir las teorías subjetivas de la acción de los productores de datos, se podrían buscar secuencias de evaluaciones de la situación, la reflexión sobre acciones alternativas, la atribución de resultados esperados y la evaluación de consecuencias personales como la satisfacción o la decepción.

4.3.3 Elaboración de códigos temáticos

El enfoque más radical consiste en intentar reducir la codificación de un fichero de datos a una categoría central, es decir, su mensaje o tema central. Dado que los análisis cuantitativos proceden como procesos cíclicos, no se puede asignar un lugar fijo a la reducción temática, como por ejemplo al final del análisis para resumir las conclusiones. Aunque la condensación de los datos en una categoría central cumple una función importante en este punto, el enfoque también es muy útil en otras colocaciones:

1. En el caso de archivos relativamente cortos y homogéneos, la reducción temática puede estar al principio del proceso de análisis de datos.
2. En el caso de archivos largos, heterogéneos y confusos, la reducción de datos puede iniciarse con el enfoque temático. Cuando se utiliza de este modo, la reducción temática cumple importantes funciones heurísticas. En primer lugar, hay que intentar no dejarse confundir por la multitud de detalles, quizá también contradictorios, y filtrar la(s) idea(s) principal(es).
3. Por regla general, la reducción temática se produce al final de la secuencia de pasos analíticos, a menudo sólo al final de varios ciclos de reducción de datos, reconstrucción y comparación de sistemas de significados, en los que se elaboraron gradualmente los "leitmotiv" de varios archivos (cf. Shelly & Sibert 1992; Strauss & Corbin 1990).

Strauss y Corbin (1990) recomiendan cinco pasos para la reducción temática. Esta secuencia es compatible con cada una de las tres ubicaciones descritas del enfoque en el ciclo analítico y ayuda a evaluar mejor la contribución de las características del software:

- Se empieza por identificar una idea central, una categoría central;
- entonces se buscan categorías subordinadas;

- a menudo será necesario diferenciar o vincular estas categorías;
- las relaciones hipotéticas entre las subcategorías y éstas y el tema se contrastan con los datos;
- las incoherencias, por ejemplo, las categorías incoherentes, dan lugar a nuevos ciclos de análisis.

4.4 Reconstrucción de los sistemas de significado

El análisis cualitativo de construcción de la teoría consiste en generar inductivamente, a partir de los sucesos o afirmaciones de los datos, la teoría sobre el tema que está subjetivamente disponible para los autores de los datos. Hablamos aquí de la reconstrucción de sistemas de significado. El problema de cómo encontrar los vínculos subjetivos se puede intentar resolver de forma inductiva, deductiva o combinando estrategias inductivas y deductivas. El enfoque que se prefiera depende sobre todo de la pregunta de investigación.

Por ejemplo, cuando se analizan grabaciones o transcripciones de entrevistas poco estructuradas o textos de forma libre, por ejemplo, entradas de un diario, se suele empezar por identificar unidades individuales de significado y asignar los segmentos de archivo correspondientes a categorías específicas. A continuación, se buscarán las secuencias típicas de estas categorías. Por último, se intenta subordinar dichas secuencias a categorías más abstractas, al tema o temas del orador. Se empieza de forma inductiva, pero se pasa de las conclusiones inductivas a las deductivas como muy tarde en la reducción temática.

Si uno se acerca a los datos con un interés cognitivo específico o una orientación teórica determinada, buscará ciertas conexiones desde el principio. Así, se parte de hipótesis sobre posibles correlaciones y se intenta demostrarlas de forma deductiva a partir de los datos. Las discrepancias y los fallos dan lugar a secciones inductivas de análisis, con las que de nuevo se puede modificar el sistema hipotético de orientación.

En el procedimiento inductivo, se intenta generalizar las categorías y sus relaciones sistemáticas a partir de los datos. En el procedimiento deductivo, se buscan segmentos de archivo específicos que puedan confirmar las presuposiciones o hipótesis generales sobre la conexión de las categorías. La siguiente visión general de las técnicas disponibles en AQUAD para la reconstrucción asistida por ordenador de las relaciones sistemáticas está estructurada según este esquema.

4.4.1 Reconstrucción de las condiciones simples

Según el avance del proceso de análisis, se pueden distinguir de nuevo dos variantes:

(1) Buscar secuencias simples de significados. Para la búsqueda, se determina un rango de datos antes y después de los segmentos de datos de una categoría seleccionada, es decir, se determina cuántas líneas o unidades de contador (fotogramas en el caso de los vídeos, décimas de segundo en el caso de los audios) deben buscarse antes y después de los segmentos de datos de la codificación seleccionada. Dentro de este rango de datos, se debe registrar la aparición de otras codificaciones. Si ciertas codificaciones se producen con frecuencia en las proximidades del código crítico, se puede comprobar una posible sistemática detrás de estas combinaciones. En AQUAD, la opción "Búsqueda" de "Secuencias de Codificación" en el módulo "Búsqueda" admite este enfoque de reconstrucción.

(2) Búsqueda de relaciones de significado condicional. La "condición" para la posible existencia de correlaciones sistemáticas está definida por (al menos) una segunda categoría (en al menos dos variantes de significado, por ejemplo, "género") a la que debe asignarse simultáneamente un segmento de datos. Para la búsqueda asistida por ordenador,

se construyen matrices de codificación (Miles & Huberman, 1994) o tablas de codificación (Shelly & Sibert, 1992). La doble determinación del contenido (segmentos de datos) de las celdas de una tabla de este tipo determina la interpretación de los resultados con más fuerza que la mera constatación de una acumulación de proximidad espacio-temporal de categorías inicialmente independientes. Por otro lado, la construcción de una matriz de análisis requiere un mayor trabajo conceptual previo, es decir, un mayor avance en el proceso de análisis. El módulo "Tablas" está disponible en AQUAD para este enfoque de reconstrucción.

4.4.2 Reconstrucción de las relaciones condicionales

También en este caso se pueden distinguir dos variantes, que corresponden a las estrategias simples y complejas de codificación secuencial descritas anteriormente:

(1) Comprobación de secuencias de codificación simples. Supongamos que de una entrevista a los padres sobre las prácticas de crianza de la familia surge la impresión de que un padre hace un esfuerzo especial para justificar sus medidas, entonces se podría utilizar el módulo "vínculos" para buscar específicamente secuencias de categorías relevantes, por ejemplo, para la codificación final o causal, y así comprobar la suposición.

(2) Comprobación de patrones de codificación complejos. Para comprobar las relaciones de codificación que superan el grado de complejidad de las estructuras de vinculación que se ofrecen en el módulo "Vinculaciones", en Aquad se dispone de un componente de programa "abierto" para formular vinculaciones de codificaciones complejas y, por tanto, específicas para cada caso.

4.5 Comparación de contextos de significado

Las comparaciones constantes de las interpretaciones dentro de un conjunto de datos (texto, vídeo, grabación de audio, imagen) y entre diferentes conjuntos de datos constituyen el núcleo de los procedimientos de análisis cualitativo (cf. Shelly & Sibert 1992). Ya en el caso de la reducción de datos dentro de un fichero, no se puede lograr una codificación fiable sin comparaciones de las categorías o los segmentos de datos tanto entre sí como entre los demás ficheros. Sin embargo, en la mayoría de los estudios se quiere captar la singularidad de cada conjunto de datos y las opiniones, competencias, visiones subjetivas del mundo, etc. que se expresan en él, más que con un sistema de categorías que sea válido para todos los archivos. En algún momento del proceso de investigación, uno quiere elaborar configuraciones generales a través de los archivos. Por lo tanto, en los análisis cualitativos hay que lidiar con el peligro proverbial de que a menudo no se pueda ver el bosque por los árboles. Por encima de las diferencias y su especificación, no hay que perder de vista lo que es común. Para ello, hay que reconocer y comparar las conexiones sistemáticas entre los archivos.

Sin embargo, el problema de comparar las correlaciones en los fenómenos sociales es que a menudo se dan multitud de condiciones en combinaciones diferentes, a veces incluso contradictorias. En muchos estudios, uno se enfrenta en última instancia a la tarea adicional de tener que comparar sus propios resultados con otros estudios, es decir, básicamente hay que realizar un meta-análisis cualitativo. Aunque otros investigadores hayan trabajado en cuestiones comparables, sólo habrán encontrado respuestas parcialmente coherentes. En la búsqueda de relaciones condicionales para un determinado fenómeno, se demuestran las más diversas constelaciones a través de diferentes estudios en todas las ciencias que se ocupan de la complejidad de los sistemas naturales.

Tomemos como ejemplo una pregunta a menudo debatida amargamente por los opositores en la vida cotidiana: "¿Existe una conexión entre el cáncer de pulmón y el tabaquismo?" Los que quieren responder afirmativamente a esta pregunta se enfrentan regularmente a la referencia a su abuelo de 80 años que había disfrutado de la bruma azul desde su temprana juventud o a la referencia a una víctima de cáncer de pulmón que nunca había fumado. Obviamente, hay muchas condiciones que influyen. Los estudios empíricos proporcionan una variedad de argumentos empíricos, de los cuales uno puede seleccionar los apropiados para la inmunización argumentativa, dependiendo del interés personal o de la propia preocupación. Sólo una comparación de los resultados de todos los estudios pertinentes o, al menos, de una muestra representativa de los mismos puede aportar claridad sobre las constelaciones de condiciones pertinentes.

Ragin (1987) ha desarrollado un procedimiento comparativo explícito aplicando el álgebra de Boole a los datos cualitativos. El procedimiento se basa en el algoritmo Quine-McClusky de "minimización lógica". Este procedimiento cumple en gran medida los requisitos del método comparativo buscado, al menos la aplicación del álgebra de Boole crea las condiciones ideales. Según Ragin (1987, p. 121), puede utilizarse para cumplir los requisitos

- para comparar un gran número de textos o casos individuales;
- para captar los vínculos complejos entre las condiciones;
- si se desea, para producir reconstrucciones "parsimoniosas";
- examinar los textos individuales tanto en sus partes relevantes como en su conjunto (véase la reducción temática);
- para comparar las reconstrucciones que compiten entre sí.

Si comparamos este procedimiento con los enfoques orientados a las variables y que parten del supuesto básico de la agregación aditiva de las variables individuales, podemos señalar que la comparación orientada a los casos está diseñada para ello (cf. Ragin 1987, p. 51 y ss.),

- elaborar invariantes o conexiones constantes de significado a través de cuidadosas comparaciones de casos individuales;
- prestar más atención a la variabilidad de las configuraciones significativas de las condiciones que a las meras distribuciones de frecuencias de los casos típicos, de lo que se deduce que hay que tener en cuenta incluso un solo caso contradictorio;
- captar los casos como conjuntos, es decir, ver las condiciones del caso en interdependencia, lo que constituye este caso particular, pero no las condiciones individuales en dependencia de una distribución de la población;
- para examinar con precisión cómo las condiciones generales inferidas conducen a resultados diferentes en distintos contextos y en distintas configuraciones.

Para poder aplicar las reglas del álgebra de Boole a las interpretaciones de los datos, se reducen radicalmente los significados encontrados en cada archivo a valores de verdad. Por lo tanto, uno se contenta con la notación binaria "la condición es verdadera" (es decir, dada) o "la condición es falsa" (es decir, no dada). No importa si se trata de condiciones genuinamente cualitativas, es decir, de interpretaciones y asignación de categorías adecuadas, o de valores medidos de una característica cuantitativa.

AQUAD ofrece las posibilidades de minimización lógica en el módulo "Implicantes". Si las operaciones comparativas se hacen necesarias en el proceso de análisis cualitativo -y esto puede ser en todas las fases debido a los ciclos de análisis- también hay que pensar en las posibilidades de minimización lógica.

Como heurística, ya rinde bien en la generación de categorías para la interpretación, aunque sólo se hayan analizado unos pocos archivos, por lo que la base de comparación sigue siendo estrecha. Hacia el final del análisis, al resumir los resultados o cuando se quiere agrupar los hallazgos individuales, distinguir los tipos de archivos o hablantes, destacar los textos clave, etc., el procedimiento de minimización lógica parece indispensable. Porque, como señala Ragin (1987, p. 51), con cada expediente adicional la carga de trabajo crece geométricamente, pero con cada categoría o condición adicional que se considera en la comparación crece exponencialmente. La agrupación o clustering de casos individuales con la ayuda de la minimización lógica fue particularmente desarrollada en AQUAD.

Por último, el procedimiento de minimización lógica puede utilizarse si se pretende comparar varios estudios cualitativos. Este es el paso metodológico al que se suele aplicar el término "meta-análisis". Cuando los resultados de la investigación se han expresado de forma cuantitativa, se puede recurrir a una serie de procedimientos muy formalizados para el meta-análisis. En el ámbito de la investigación cualitativa, el enfoque de minimización lógica puede aportar la deseable simplificación, transparencia, fiabilidad y documentabilidad de las comparaciones meta-analíticas.

Capítulo 5: El análisis de textos como análisis de secuencias

El objetivo del análisis de secuencias es establecer una hipótesis de estructura de caso mediante la formulación y comprobación de hipótesis sobre el significado de cada segmento del texto. La estructura del caso permite hacer afirmaciones generales más allá de las sensibilidades individuales de los actores, ya que éstos deben orientarse para su comunicación e interacción a un sistema lingüístico-cultural común de reglas. Las explicaciones del contenido de este capítulo y dos ejemplos están tomados de Gürtler, Studer y Scholz (2010) y Studer (1988) respectivamente.

El módulo "Métodos de análisis" -> "Hermeneútica objetivo" de AQUAD divide el trabajo con los textos en tres secciones:

- Como requisito previo al análisis secuencial, los textos deben dividirse en segmentos de texto, a los que posteriormente
- se elaboran sucesivas hipótesis sobre todos sus significados concebibles, antes de que
- se someten éstos posteriormente a un examen crítico en otros pasajes del texto.

5.1 Preparación de los textos

En el análisis secuencial, la interpretación sigue la secuencia de acontecimientos o interacciones exactamente como se registran en el texto. Wernet (2009, p. 27) subraya la importancia de la "actitud interpretativa básica", según la cual el "texto debe tomarse en serio como texto" y no hay que "evaluarlo como una cantera de información o como una feria de ofertas de sentido" ni tampoco "explotarlo". Este procedimiento está condicionado por el objetivo del análisis de secuencias, a saber, según Oevermann et al. (1979), determinar las estructuras de *significado latentes* de la interacción en los textos como protocolos "de acciones o interacciones sociales reales, simbólicamente mediadas" (p. 378) y distinguirlas del *contenido manifiesto*: "los textos de interacción constituyen estructuras de significado objetivas sobre la base de reglas reconstruibles y estas estructuras de significado objetivas representan las estructuras de *significado latentes* de la propia interacción" [cursiva en el original]. Son "realidad (y perduran) analíticamente (si no empíricamente) independientes de la respectiva representación intencional concreta de los significados de la interacción por parte de los sujetos implicados en la misma" (ibid., p. 379). Esto significa que no sabemos nada sobre las intenciones de otras personas, pero podemos reconstruir sus motivos a partir de las transcripciones de los textos.

Según esta postura, el significado general de un texto, es decir, de un protocolo de interacción, no se revela tratando de desentrañar los motivos, las intenciones de actuar, las ideas normativas, etc. de los participantes, "la realidad psíquica interior de los sujetos de la acción" (Oevermann et al., 1979, p. 379), sino intentando reconstruir su estructura objetiva de significado. Esta estructura de significado existe independientemente de los estados de ánimo, las perspectivas y las intenciones de los agentes como una "realidad social", como una "realidad de posibilidades" (Oevermann et al., 1979, p. 368), que están disponibles a través del sistema social dado de normas y reglas - o también pueden inferirse de una violación de las reglas. Por lo tanto, en la interpretación secuencial del texto es esencial distinguir claramente entre el nivel de la estructura objetiva del significado y el nivel de los significados subjetivos. Sobre todo, la estructura objetiva de significados o la estructura latente de significados de las interacciones debe examinarse independientemente y antes de tener en cuenta los significados subjetivos accesibles a los participantes. Oevermann et al. (1979, p. 380), refiriéndose a George Herbert Mead, especifican que parten de una noción de significado como una estructura social objetiva que aparece en la interacción - y "que a su vez debe considerarse como una condición previa para la intencionalidad [de los participantes; adición del autor]". Por otra parte, esto significa que la autocomprensión de los sujetos que interactúan, sus intenciones y motivos, ciertamente juegan un papel en la interpretación, pero sólo en el contexto de la estructura latente de significado de su interacción, es decir, después del análisis.

Oevermann et al. (1979, p. 354) consideran este enfoque "una perspectiva muy simple, que argumenta casi con supuestos triviales". Sin embargo, de estos supuestos se derivan estrictas consecuencias metodológicas. En particular, hay que tener en cuenta "que no se puede utilizar la información y las observaciones de interacciones posteriores para interpretar una interacción precedente" (Oevermann et al., 1979, p. 414). O, como señala explícitamente Wernet (2009, p. 28): "Uno no deambula por el texto en busca de pasajes útiles, sino que sigue el protocolo textual paso a paso." Y además: "Para el análisis de la secuencia es sumamente importante *no prestar atención* al texto que sigue a un pasaje que debe interpretarse" [cursiva en el original]. Esto no implica que se deba interpretar un texto desde la primera frase o que no se deban omitir pasajes, sino que los pasajes seleccionados como relevantes para la pregunta de investigación sólo deben interpretarse secuencialmente. Para una mejor comprensión, una repetición: el objetivo del análisis secuencial es examinar una estructura que se desarrolla de forma natural exactamente de la misma manera; para ello, sin embargo, el orden cronológico de los acontecimientos en el protocolo del texto es fundamental.

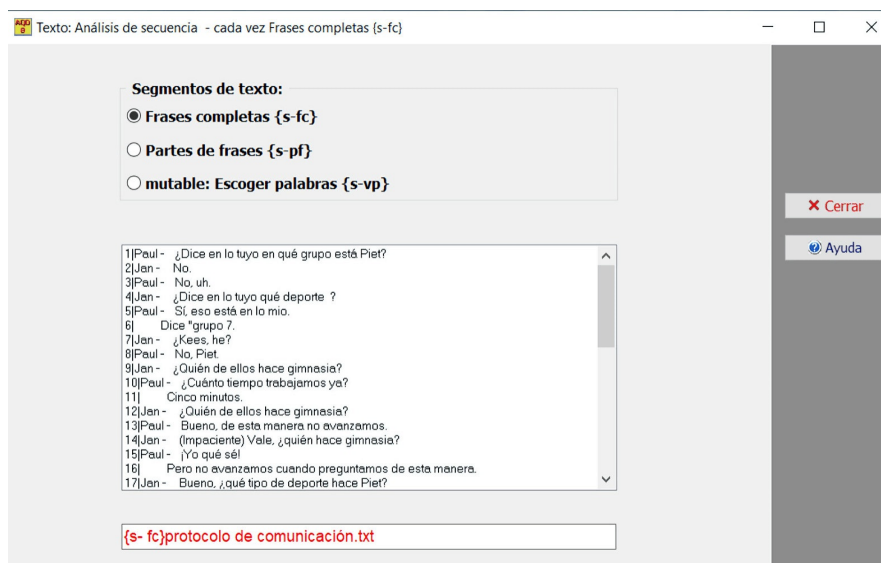
Por lo tanto, como primer paso, el texto debe dividirse en unidades de secuencia, en segmentos de texto sucesivos, para el trabajo de interpretación en el ordenador. En AQUAD se puede

- fragmentar el texto según las partes de la frase (hasta el siguiente signo de puntuación (.,:;!)) o
- trabajar con frases completas.

También es posible

- dividir el texto en segmentos de texto palabra por palabra.

En el siguiente ejemplo, el texto se ha dividido en frases completas. En el renombramiento de los textos estructurados secuencialmente, el nombre del texto original (ejemplo en AQUAD) "Protocolo de cooperación.txt" va precedido por "{s-fc}", lo que pretende señalar que se trata de un texto estructurado secuencialmente según frases completas: "{s-fc}Protocolo de comunicación.txt". En consecuencia, "{s-pf}" representa la estructura según las partes de la frase, "{s-vp}" para los segmentos de texto compuestos de forma variable según diferentes números de palabras (vp - varias palabras).



5.2 Generación de hipótesis

5.2.1 Antecedentes teóricos

"El objeto concreto de los procedimientos de la hermenéutica objetiva son los protocolos de las acciones o interacciones sociales reales, simbólicamente mediadas, ya sean escritas, acústicas, visuales, combinadas en diversos medios de comunicación o archivadas de otro modo" (Oevermann et al. 1979, p. 378). El significado latente se encuentra en estos protocolos: este es el objetivo de la hermenéutica objetiva. Arriba ya se esbozó que para ello, según Wernet (2009, p. 39 y ss.), se sigue una secuencia de "contar historias" (historias en las que podría ocurrir el segmento de texto), "formar lecturas" (es decir, ordenar comparativamente las historias según las similitudes y diferencias) y "confrontar estas lecturas con el contexto real" (compararlas sistemáticamente).

No sólo hay que tener en cuenta la experiencia única de las personas, sino también sus realidades objetivas, que pueden aprovecharse como patrones de acción individuales, sociales y culturalmente arraigados. La estricta orientación hacia lo que es, hacia la realidad, guía el procedimiento en la Hermenéutica Objetiva. El punto de partida es el caso general (normal), que se hace público y generalmente accesible a través de la reconstrucción de la estructura latente del significado. Después, la manifestación concreta, la forma de expresión puede ser interpretada con mayor precisión por la persona empírica y su subjetividad.

Para ilustrar esto, nos gustaría dar un ejemplo tomado de la carta de solicitud de un cliente potencial para la terapia de adicción (véase Studer, 1996). Una formulación común de la motivación de la terapia que se encuentra allí dice:

"Sigo muy interesado en controlar mi adicción".

Podemos leer este segmento de texto como "quiero liberarme de la adicción" o "quiero ser capaz de vivir de forma abstinente", pero también "quiero ser capaz de controlar mi consumo de drogas" - porque lo que tienes agarrado, no lo sueltas. Un análisis más detallado de la carta de solicitud (véase Studer, 1996) muestra que, en el caso de este cliente potencial, la motivación inicial "para controlar la adicción" se dirige al control, pero no a la superación

completa de la adicción. El control debe evitar, en la medida de lo posible, que el cliente sufra daños. El conocimiento de las estructuras de significado latentes reconstruidas de este modo aporta la inestimable ventaja para la práctica de estar cerca de la realidad de los clientes, que repercute en sus acciones y determina su vida cotidiana. En primer lugar, hay que tomar en serio y valorar su motivación expresada en el momento de cambiar su modo de vida anterior. Sin embargo, a esto le sigue una exploración realista del marco de posibilidades para poder abandonar a tiempo las ilusiones terapéuticas y adaptar las ofertas terapéuticas individualmente. Para ilustrar esto, nos gustaría citar otro segmento de texto de otra carta de un cliente (véase Studer, 1996):

"También, sin embargo, nosotros [mi pareja y yo] tuvimos muy buenas conversaciones ..."

Podemos leer este segmento como otro ejemplo ("También tuvimos...") del aprecio de la escritora por su pareja. Sin embargo, también podemos leer ("Pero también tuvimos...") que la escritora tiene dos pensamientos diferentes en su cabeza al mismo tiempo: "Por un lado mi pareja me molestaba, pero por otro lado también podíamos tener una buena conversación". Sin embargo, como muestra el análisis posterior de la estructura latente, la persona que escribe no mantiene claramente separados estos dos pensamientos y mezcla estas intenciones de acción separadas de acercamiento y distanciamiento. En el trabajo práctico con el cliente, la separación de los dos pensamientos abre a su vez el espacio de posibilidades para descubrir cosas nuevas sobre uno mismo. Pero primero hay que entender qué es lo que realmente está en juego en ese momento antes de comenzar las intervenciones terapéuticas. ¿Qué realidad procesable se esconde detrás de las expresiones verbales observadas? ¿Qué motiva la acción de un momento a otro? Es imprescindible abordar estas cuestiones con éxito para poder actuar con flexibilidad en la terapia y ampliar el espacio de posibilidades en lugar de estrecharlo aún más. Lo que los clientes hagan con esta interpretación es otra cuestión.

Oevermann (2002, p. 33) afirma: "El análisis de secuencias anida en la estructura básica de los acontecimientos humano-sociales reales y, por lo tanto, no es, como los procedimientos habituales de medición y clasificación, un método externo al objeto, sino uno correspondiente y apropiado para el objeto mismo. De hecho, en la vida práctica, las decisiones también deben tomarse en principio en cada punto de la secuencia entre las opciones que aún están abiertas a un futuro abierto." Análisis de la secuencia "... se apoya en la secuencialidad que es constitutiva de la acción humana" (ibíd., p. 6). Sin embargo, un enfoque secuencial no significa simplemente trabajar de adelante hacia atrás. Se trata más bien de abrir nuevas posibilidades (se utiliza el término lectura) estrictamente en la secuencia del texto a analizar y cerrarlas de nuevo cuando no resisten el escrutinio del texto. El análisis de la secuencia es, por tanto, la interacción entre la posibilidad (de "opciones abiertas en un futuro abierto", véase más arriba) y la realidad (o el intento de volver a encontrar estas opciones en el texto).

5.2.2 Principios de generación de hipótesis

Como ya se ha citado, Wernet (2009, p. 39) da una sencilla "respuesta a la pregunta: ¿qué tengo que hacer para llevar a cabo una operación metodológicamente verificable de reconstrucción del significado según reglas válidas?" La respuesta consiste en un proceso metodológico de tres pasos, de los cuales los dos primeros nos interesan aquí para la generación de hipótesis: "Tengo que (1) contar historias, (2) formar lecturas...". El tercer paso, la confrontación de las lecturas con el contexto real, nos ocupará en la comprobación de las hipótesis (véase el capítulo 5.2.3).

Antes de empezar a interpretar los segmentos de texto, hay que aclarar cuál es el caso y en qué contexto se inscribe. Según Wernet (2009), esto incluye aclarar y revelar el interés de la investigación, por un lado, y por otro, aclarar qué es lo que el protocolo de texto realmente registra, qué realidad social registra o, más concretamente, qué contribución puede hacer una entrevista, por ejemplo, para responder a las preguntas de la investigación. Comenzamos

con los datos objetivos del caso (nacimiento, entorno, ocupaciones, fechas de fallecimiento, etc.), y sólo entonces sometemos al análisis el texto propiamente dicho (conversaciones en el aula, conversaciones terapéuticas, entrevistas biográficas, etc.). La recogida de datos objetivos corresponde, por ejemplo, a la elaboración de un genograma (datos familiares).

La interpretación se guía por las cinco reglas de independencia del contexto, literalidad, secuencialidad, extensividad o totalidad y parsimonia (véase Wernet, 2009, pp. 21-38):

Independencia del contexto:

Por supuesto, una interpretación textual debe tener en cuenta el contexto en el que surgió el protocolo, pero sólo después de que se hayan desarrollado los posibles significados de un segmento del texto independientemente de ese contexto. Así, primero se inventan historias en las que el pasaje del texto crítico podría tener sentido. Wernet habla aquí de diseñar exclusivamente "contextos de experimentación del pensamiento" para generar posibles significados del segmento de texto en el primer acceso al texto. Wernet advierte que la interpretación dependería de la comprensión cotidiana del intérprete o de la comprensión previa no obtenida científicamente, por lo que podría ponerse en marcha un proceso circular de interpretación si no se deja de lado metódicamente la comprensión previa por el momento.

Literalidad:

La literalidad significa que la interpretación se guía precisa y exclusivamente por lo que ocurre real y verificablemente en el texto, y no por lo que el texto podría haber querido decir. El texto como tal es una pieza de la realidad y debe ser tratado como tal. De ello se derivan dos condiciones importantes:

- (1) El protocolo del texto debe basarse -en realidad, como es lógico- en la interpretación en su forma original, por ejemplo, como transcripción palabra por palabra de una interacción social, y no en forma de paráfrasis alienada por el entendimiento cotidiano y las convenciones sociales, al menos reducida en sus posibilidades de significado de antemano. Estas paráfrasis, que a veces se utilizan en los análisis de contenido cualitativos, están absolutamente fuera de lugar.
- (2) La interpretación debe ser francamente exigente con el texto. No debe pasar por alto los detalles que parezcan poco adecuados o incluso inapropiados, como podríamos hacer en la vida cotidiana. Wernet aboga por "poner el texto 'en la escala de oro' de una manera ..." que en la vida cotidiana "... parecería insignificante".

De este modo, inapropiado en las conversaciones cotidianas, quizá incluso escandalosamente mezquino, se revela el significado latente tras la formulación manifiesta. En los ejemplos anteriores, "Quiero controlar mi adicción" y "Pero también hemos tenido muy buenas conversaciones", la interpretación literal choca con la desviación de lo que se quiere decir con respecto a lo que realmente se dice y, por tanto, con la ambigüedad de estas afirmaciones y las discrepancias en el mundo vital de la persona.

Secuencialidad:

Este principio determina el núcleo del análisis, ya que el procedimiento de interpretación es estrictamente lógico y está orientado al paso a paso. Esto hace que el procedimiento sea científico. Así, no se buscan pruebas de posibles interpretaciones (lecturas) de forma transversal y no sistemática a través del texto (búsqueda de confirmación positiva).

Más bien, se practica una alternancia de formulación de hipótesis, condensación de las mismas en lecturas y comprobación con el texto en estricta secuencia.

Dónde comienza exactamente la interpretación, qué segmento de texto representa la secuencia de apertura, debe justificarse en términos de contenido. No es necesario comenzar con la primera frase o parte de una frase. Pero entonces uno sigue metódicamente la secuencia del protocolo de forma estricta, paso a paso. Lo que sigue en el texto después del segmento que se está enfocando en ese momento debe - en un principio - ser ignorado. Por supuesto, cada segmento de texto está integrado en el "contexto interno" del significado del texto desarrollado hasta ahora. El desconocimiento de los respectivos segmentos de texto siguientes es altamente significativo metodológicamente y prácticamente: "La progresión pensamiento -experimental deja claro que el respectivo caso concreto debe tomar una 'decisión' de ser lo que es...". Wernet subraya que esto "...apunta a la reconstrucción del 'devenir-así-y-no-otro' de una práctica vital".

El principio de secuencialidad no prohíbe saltarse segmentos del texto y seleccionar pasajes relevantes del mismo en función de la pregunta de investigación y del progreso de la reconstrucción de la estructura de significado del texto. Sin embargo, su inicio debe justificarse de nuevo, la interpretación debe proceder de nuevo secuencialmente paso a paso. La selección de los pasajes del texto posterior se basa en si pueden confrontar o incluso apoyar las diferentes lecturas.

La secuencialidad también significa no dejar que el conocimiento previo sobre los desarrollos posteriores del texto fluya en la interpretación, sino orientarse al desarrollo paso a paso al interpretar. Es importante no sólo buscar la confirmación de las hipótesis, sino precisamente las contrapruebas para poder invalidarlas. Esto corresponde al enfoque falsacionista del racionalismo crítico (Popper, 1943). La reconstrucción del significado debe ceñirse al texto y verificarse exactamente en él. En la práctica, sin embargo, se suelen encontrar intentos de apoyar empíricamente las propias ideas del modelo en lugar de cuestionarlas. Por lo tanto, el examen justo de las hipótesis que compiten entre sí sólo tiene lugar de forma limitada.

Extensividad:

La extensividad o totalidad es un principio de equilibrio que debe evitar que se sigan intenciones subjetivas y arbitrarias a la hora de interpretar. Así pues, no hay que pasar por alto detalles aparentemente sin importancia, pero tampoco hay que limitar la interpretación a los pasajes aparentemente importantes. "Importante" y "no importante" parecen ser el resultado de prejuicios subjetivos. Más bien, es importante proceder a la interpretación de manera que el texto se analice paso a paso sin favorecer o desfavorecer nada. Es esencial tratar sistemáticamente todos los pasajes por igual. En particular, "... los contextos de pensamiento-experimentación también deben ser plenamente iluminados tipológicamente...".

Parsimonia:

Este principio exige que se asuma la normalidad, entendida como lo cotidiano con sus rutinas, a la hora de construir posibles interpretaciones de la práctica vital. La demanda aquí es "... admitir sólo aquellas lecturas que sean textualmente verificables". Esto restringe la "narración" a las variantes que son compatibles con el texto en su conjunto. Así, el principio de parsimonia en la interpretación de la práctica vital excluye los relatos de ciencia ficción, las divagaciones esotéricas o, sobre todo, las atribuciones infundadas de una desviación patológica del comportamiento normal.

Una desviación de la norma, que a menudo también se considera patológica, es muy fácil de notar en los protocolos, porque aquí el margen de interpretación es drásticamente limitado. Las rutinas y las soluciones de problemas sostenibles son menos llamativas porque el espectro de posibilidades es muy amplio. Por tanto, la normalidad es más difícil de reconstruir que la desviación de la norma. Por lo tanto, "ahorrar" debe entenderse en dos sentidos: Por un lado, la normalidad se da inicialmente por sentada y la asunción de la desviación está sujeta a la obligación de justificar. Se trata de descubrir desviaciones realmente significativas de la normalidad de la acción humana. Por otra parte, el espacio de posibilidades se ve limitado por el hecho de que la interpretación sólo se realiza sobre el texto existente y todas las conclusiones se basan únicamente en el texto, no en las experiencias cotidianas.

En resumen, la máxima es: lo que se interpreta debe fundamentarse en el texto y lo que se fundamenta pertenece a la interpretación.

5.2.3 Condensar las hipótesis en lecturas del texto (agrupación)

La totalidad de las reglas de secuenciación o de generación de significados, que generan una multiplicidad de opciones de nuevo en cada punto de la secuencia, en cada segmento del texto, es un conjunto de reglas algorítmicamente diferentes. Como ejemplos de estas reglas, Oevermann (1996b, p. 7) menciona la sintaxis lingüística, las reglas pragmáticas de la acción del habla, así como las reglas lógicas del razonamiento formal y material-sustantivo (inducción, abducción, deducción; cf. Reichertz, 2000). En cuanto a las interpretaciones que surgen de las posibilidades, Oevermann habla de enunciados bien formados. Existe un enunciado bien formado cuando se ha generado una interpretación según las reglas y con los bloques de construcción de una lengua, cuando se trata de un enunciado lingüístico normal, escrito u oral, que se ajusta a todas las reglas de la lengua respectiva en términos de elección de palabras, estilo y gramática. Se trata, pues, de un enunciado gramaticalmente correcto y, por tanto, de *una frase completamente normal*.

Las hipótesis se agrupan para derivar tipos. En su formulación, estos tipos forman las lecturas del texto. Sólo después de este paso se prueban las lecturas sobre el texto y se confrontan con la realidad del mismo. De este modo, la conexión entre lo general (estructura abstracta e independiente del caso) y lo particular (práctica vital concreta) se perfila con mayor claridad y precisión.

5.3 Verificación de las hipótesis

Una vez reconstruida la estructura básica, se comprueban las hipótesis generadas en el proceso. Técnicamente, se aplica el principio de falsificación (Popper, 1934). Esto significa que una hipótesis no se mantiene porque se pueda demostrar, sino que se mantiene porque su contrario, su inaplicabilidad al caso concreto, no se ha podido demostrar. Por lo tanto, no existe una verdad siempre válida, sino sólo hipótesis cuya no-verosimilitud aún no ha sido demostrada. Probar una hipótesis equivale, por tanto, a intentar demostrar y justificar que la hipótesis no es compatible con el texto.

Esto ya no tiene que hacerse de forma secuencial. En la búsqueda de pruebas para los relatos generados en el curso de la generación de hipótesis o las lecturas del texto diseñadas a partir de ellas, es necesario "divagar" en el texto, concretamente para encontrar precisamente esas contrapruebas.

Las posibles dificultades en la verificación se derivan sobre todo de la violación de las reglas de interpretación del texto. Por lo tanto, las conjeturas audaces pueden no haber sido expresadas y realmente perseguidas, aunque es precisamente la posibilidad potencial de que una conjetura fracase lo que constituye su poder explicativo. Una

hipótesis es sólida si asume el riesgo de ser derribada incluso por una pequeña discrepancia. De ello se deriva un mayor ámbito de validez. Una hipótesis que no puede ser refutada en principio carece de valor científico.

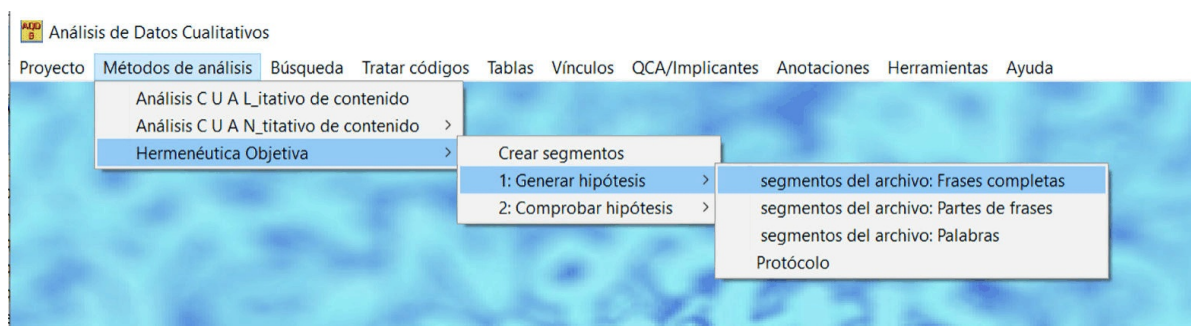
Por tanto, al recurrir a las recomendaciones sobre la generación de hipótesis, hay que añadir que hay que investigar todas las posibilidades, por pequeñas que sean, de cómo se debe entender un texto. Naturalmente, este requisito hace que el procedimiento sea muy laborioso, pero también lleva a que una muestra muy pequeña sea suficiente para cubrir exhaustivamente un campo de investigación. El propio Oevermann llega a afirmar que se pueden hacer afirmaciones generales con un solo caso bien analizado. Hildenbrand (2006) parte de la base de que son necesarios entre ocho y diez casos para alcanzar la saturación teórica en el sentido de la estrategia de investigación de la teoría fundamentada/ancorada según Glaser (1998); después de eso, los casos adicionales no prometen ninguna ganancia sustancial adicional de conocimiento.

El análisis continúa hasta que la estructura del caso se reproduce completamente una vez (Oevermann, 1996). Dado que la lógica interna se reproduce y se desarrolla de nuevo en cada sección, basta con examinar algunos pasajes con gran detalle para seguir obteniendo una visión de conjunto del caso. Esta estabilidad de la interpretación textual legitima el análisis (sólo) hasta la (primera) repetición de la estructura y la búsqueda de pruebas contrarias para invalidarla. Esto se ve acentuado por el hecho de que uno "sólo [realmente] conoce una ley de estructura de casos o una estructura de casos... sólo [se sabe realmente] cuando se ha reconstruido una fase completa de su reproducción o transformación mediante el análisis de la secuencia" (ibíd., p. 9). Del análisis surgen entonces las estructuras de sentido manifiestas y evidentes o latentes de la rutina diaria en situaciones estándar (ibíd., p. 76 y ss.).

5.4 Generación y comprobación de hipótesis de análisis de secuencias con AQUAD

La preparación de los textos ya se ha descrito anteriormente. A continuación, nos adentraremos en la generación y verificación de hipótesis utilizando el ejemplo del fragmento de texto "Cooperation protocol.txt". Este texto se guardó como ejemplo de proyecto "Comunicación" en el directorio principal de AQUAD (por ejemplo, "C:\Aquad_8") en el disco duro durante la instalación de AQUAD. Ahora tenemos este texto estructurado (con la función "Preparación de textos") en frases completas. Estos están disponibles para su análisis en el subdirectorio "C:\Aquad_8\cod" como el archivo "{s-fc}protocolo de cooperación.txt".

5.4.1 Generación de hipótesis



Para empezar, seleccionamos "Hermenéutica objetiva" -> "1: Generar hipótesis" -> "Segmentos de archivo: frases completas" en el menú principal.

En primer lugar, se abre una ventana en la que se selecciona un archivo del proyecto actual para editarlo. Ahora, el proyecto "Comunicación" sólo incluye el archivo "{s-fc}protocolo de cooperación.txt". Lo seleccionamos haciendo clic sobre él.

Se abre la ventana para introducir las hipótesis. Inicialmente sólo muestra el primer segmento de texto en el campo de texto enmarcado en azul. Todo el protocolo procede de un estudio sobre cooperación realizado por Van der Linden, Erkens y Nieuwenhuysen (1995), en el que dos estudiantes tenían cada uno una carta de unos chicos en un campamento de verano y tenían que averiguar de forma cooperativa quién estaba con quién en el campamento:

1 | Paul - ¿Dice en lo tuyo en qué grupo está Piet?

Text: Análisis de secuencia 1: Generar hipótesis / Frases completas - {s-fc}protocolo de comunicación.txt

Fila	Texto
1	Paul - ¿Dice en lo tuyo en qué grupo está Piet?

M	desde	hasta	Hipótesis

desde 1 hasta 1

Cancelar Vaciar Guardar

Próx. segmento
Cerrar
Ayuda

La ventana de la derecha enumera las hipótesis que generamos una a una. Por supuesto, esta ventana todavía está vacía, ya que estamos empezando con el análisis de la secuencia. Para empezar, hacemos clic en el primer segmento de texto marcado en azul a la izquierda, tras lo cual se abre una ventana de entrada a la derecha en el panel de hipótesis para nuestra primera idea, la primera hipótesis para el segmento de texto marcado:

Text: Análisis de secuencia 1: Generar hipótesis / Frases completas - {s-fc}protocolo de comunicación.txt

Fila	Texto
1	Paul - ¿Dice en lo tuyo en qué grupo está Piet?

M	desde	hasta	Hipótesis

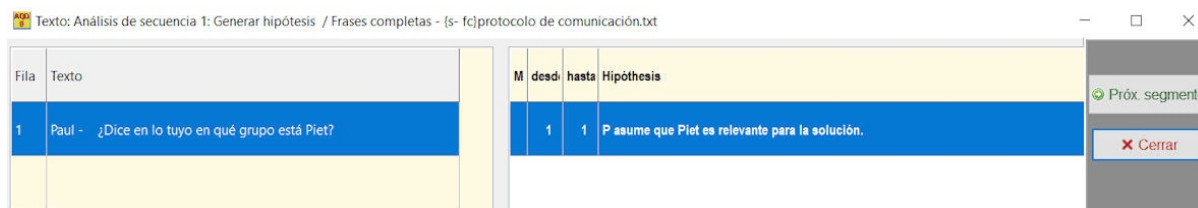
desde 1 hasta 1

P asume que Piet es relevante para la solución.

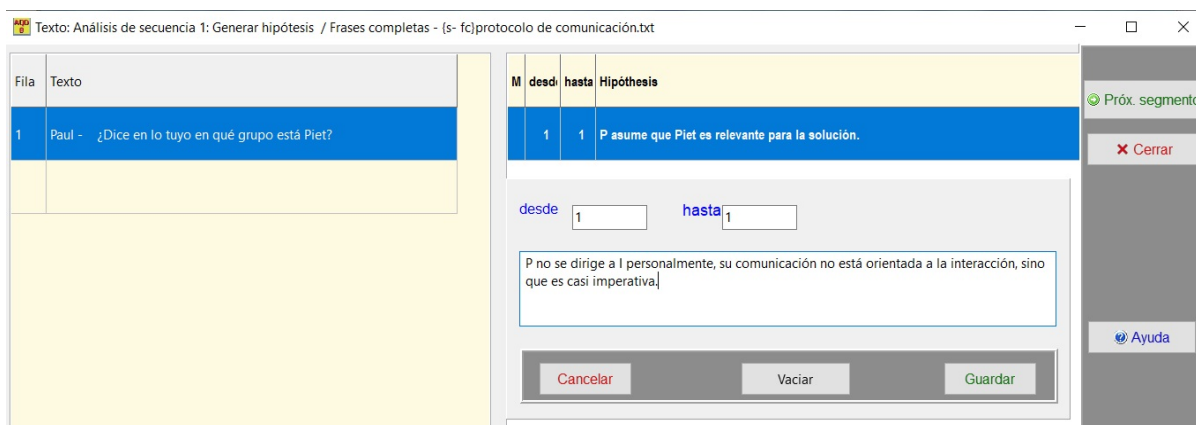
Cancelar Vaciar Guardar

Próx. segmento
Cerrar
Ayuda

Aquí introducimos nuestra hipótesis (con un máximo de 264 caracteres) y luego hacemos clic en el botón "Guardar", que traslada nuestra hipótesis a la ventana del protocolo.



¿Qué más pensamos o notamos en este segmento? Si volvemos a hacer clic en el texto de la izquierda, la ventana de entrada se abre de nuevo a la derecha. Introducimos una hipótesis más y la guardamos de nuevo:



Seguimos así hasta que no tengamos más ideas sobre este segmento. A continuación, hacemos clic en el botón "Segmento" del extremo derecho y aparece el siguiente segmento de texto a la izquierda, para el que volvemos a anotar todas las hipótesis que se nos pasan por la cabeza. Seguimos así hasta que tengamos la impresión de que las hipótesis se repiten. En este punto, dejamos de "contar historias" sobre el siguiente segmento. En su lugar, avanzamos ahora en el texto hasta que surgen nuevos aspectos en el texto que son relevantes para la pregunta de investigación. A partir de este segmento, volvemos a trabajar de forma estrictamente secuencial. Una vez que hayamos cribado todo el texto de esta manera y lo hayamos analizado secuencialmente en sus apartados correspondientes, podremos pasar a la segunda fase, la comprobación de las hipótesis.

Encontrará la lista completa de hipótesis si carga el proyecto de ejemplo "Comunicación", que se instaló durante la configuración de Aquad, y activa la "Fase 1: Generación de hipótesis". A continuación, también verá que la generación de hipótesis ha terminado después del 14º segmento de texto, ya que los patrones estructurales se repiten.

5.4.2 Prueba de hipótesis

Si hacemos clic en el número de una hipótesis (columnas desde o hasta) en la ventana de hipótesis de la derecha, se abre una ventana en la ventana de segmento de texto de la izquierda para seleccionar otras actividades. Ahí podemos aceptar la hipótesis seleccionada (en el ejemplo la hipótesis 1)

- como verdadero. Entonces el programa marca la hipótesis como verdadero con el código "T".

- o recházalo como falso. Entonces se marca la hipótesis como falso con el código "-F".

- o bien asignarlo primero a una "lectura", es decir, a un determinado grupo. Determinamos de qué grupo se trata y cómo debe nombrarse introduciendo una designación adecuada en la ventana de entrada. Esta asignación también se antepone a la hipótesis como un código ("-G-:").

La posibilidad de agrupar también puede utilizarse para identificar conexiones entre los segmentos del texto, por ejemplo, secuencias específicas, causas y consecuencias, etc. En la versión 8 de AQUAD, se incorporaron consultas específicas para este fin, por ejemplo, por "tipo" (de hipótesis) o "referencia" (a otras hipótesis). La asignación a un grupo específico "-G-:" sustituye estas clasificaciones rígidas. Detrás del marcador "-G-:" podemos anotar cualquier tipificación y conexión que tenga sentido en el proyecto individual. La función de búsqueda "Palabra clave" ayuda a encontrar de nuevo las diferentes agrupaciones y a guardarlas o imprimirlas de forma resumida.

Ahora podemos ordenar las hipótesis en función de las "lecturas" del texto asignando las hipótesis similares a una hipótesis común y ahora también designando la lectura asignándola a un grupo. Podríamos asignar la hipótesis 2 del segmento 1: "Pregunta específica y estrecha (minimalista)" (véase la figura anterior) a las hipótesis 13 (segmento 4), 23 (segmento 7), 29 (segmento 9), 40 (segmento 12) y 46 (segmento 14). Se introduciría "Pregunta estrecha" como tipo de lectura/grupación.

De nuevo: con la función de búsqueda "Palabra clave" (en el margen derecho) podemos obtener una visión general, por ejemplo, de todas las hipótesis "verdaderas" o "falsas" o de determinadas agrupaciones. Por supuesto, estos resúmenes pueden guardarse e imprimirse...

¿Y cuál es el resultado del análisis de la secuencia del proyecto de ejemplo "Comunicación"? Por qué no ejecuta la prueba de hipótesis con el material instalado o borra los archivos "{s-fc}protocolo de cooperación.shg" y "{s-fc}protocolo de cooperación.shc" en el subdirectorio ..\cod o copia estos archivos en otro directorio. A continuación, puede probar el análisis de la secuencia desde el principio con sus propias hipótesis. En cualquier caso, el resultado será una estructura de casos que ilustra cómo fracasa la cooperación cuando los socios no ponen en común sus recursos, sino que uno de ellos persigue sus propias ideas y quiere utilizar al otro sólo como proveedor de información.

Capítulo 6: Preguntas específicas de codificación

6.1 Creación de un proyecto

Los archivos de texto a codificar deben ser legibles por AQUAD, es decir, los archivos de texto deben importarse como archivos de texto plano (*.txt). A partir de entonces, AQUAD ya no trabaja con las transcripciones originales, sino con copias internas. Ahora vamos a ver el trabajo de codificación en gran detalle, paso a paso.

Al iniciar AQUAD desde el administrador de archivos o desde el escritorio, aparece la página de título. Si primero quiere combinar textos para su análisis en un proyecto, haga clic en "Nuevo proyecto" en el módulo "Proyecto", escriba un nombre de proyecto y compile una lista de los textos que se analizarán en este proyecto. Los ajustes que Ud. defina de esta manera se activan automáticamente cuando vuelva a iniciar AQUAD, hasta que introduzca

un nuevo proyecto o abra otro que ya haya definido. Por supuesto, si quiere (re)abrir un proyecto que ya ha definido, haga clic en "Abrir proyecto" inmediatamente.

Puede familiarizarse con la codificación en AQUAD trabajando con los ejemplos de texto dados. Para darle un ejemplo de cómo proceder con AQUAD, se han instalado cuatro entrevistas a profesores principiantes (en escuelas españolas) como proyecto de muestra "Entrevista" y un cuento de hadas del poeta danés Hans Christian Andersen como proyecto de muestra "Poeta". Por el simple hecho de que los textos son más breves, nos fijaremos aquí en el ejemplo del "Poeta".

El texto se divide en los archivos "poet001.txt", "poet002.txt" y "poet003.txt" y se almacena en el directorio que se especificó para la base de datos cuando se instaló AQUAD: Los archivos originales como base de datos se almacenan siempre en el directorio raíz de AQUAD, por ejemplo en "C:\Aquad8_c_texto".

Dado que AQUAD siempre activa el proyecto actual automáticamente al inicio, no es necesario introducir los ajustes una y otra vez. Se pasa directamente de la ventana de inicio al menú principal. Para ver o continuar trabajando en la presente codificación, abra primero el proyecto por defecto "poet.prj" en el módulo "Proyecto" con "Abrir proyecto". A continuación, seleccione la opción "Análisis de contenido cualitativo" en el módulo "Métodos de análisis".

Si el significado de una opción del menú no le resulta del todo claro, AQUAD ofrece ayuda general para problemas comunes en el menú principal (véase la última opción del menú principal; detalles más adelante) y ayuda contextual dentro de las diferentes ventanas.

6.2 ¿Cómo proceder a la codificación?

6.2.1 Tipos de códigos

Cuando se trabaja con AQUAD, conviene distinguir claramente entre seis tipos de códigos:

(1) Los *códigos conceptuales* se utilizan como "etiquetas para eventos distinguibles, sucesos y otros tipos de fenómenos" (Strauss & Corbin 1990, p. 61). Como ejemplo, tomemos "deliberar": Siempre que en un texto se tenga la impresión de que aquí se está dando un consejo a alguien o de que alguien está dando un consejo a otros, se adjuntará el código correspondiente, como "consejo". Los códigos conceptuales en AQUAD pueden tener hasta 60 letras y/o dígitos. No es necesario utilizar abreviaturas crípticas porque cada código sólo debe escribirse una vez. A continuación, se inscribe automáticamente en un registro de códigos. A partir de ahí, se puede seleccionar fácilmente una y otra vez haciendo clic en él.

(2) Los *códigos de perfil*, en su mayoría códigos sociodemográficos, caracterizan un fichero completo. Definen su perfil, por así decirlo. Pueden ser cualitativos o numéricos. Por ejemplo, la característica de un expediente puede ser que se trate de una entrevista con una "mujer", o que las notas de campo se refieran a una escuela "con más de 1000 alumnos". Dado que cada una de estas características sólo se necesita una vez para identificar un archivo de datos – el "género" como categoría del perfil del entrevistado no puede ser tanto "masculino" como "femenino", el "tamaño del centro educativo" no puede ser tanto inferior como superior a 1000 alumnos –, estos códigos sólo pueden aparecer una vez en cada documento de datos. Por lo tanto, también se habla de códigos singulares. Para distinguirlos de los códigos conceptuales, se ha acordado que los códigos de perfil comiencen siempre con una barra "/", por lo que en nuestros ejemplos introduciríamos:

/género: w
/tamaño de la escuela > 1000

(3) Los *códigos numéricos* se utilizan cuando el fenómeno que interesa al investigador existe como valor cuantitativo, como las puntuaciones de los exámenes, los marcadores de tiempo en la grabación de una entrevista, la edad de los entrevistados, etc. Los códigos numéricos son especialmente útiles cuando se comprueban hipótesis o dentro del análisis tabular o matricial. Sin embargo, AQUAD debe ser capaz de distinguir los códigos numéricos de otros tipos de códigos. Dado que los códigos numéricos a menudo también identifican el perfil de un archivo, por ejemplo cuando registramos la edad de los entrevistados o indicamos el tamaño de la escuela en el ejemplo anterior, se acuerda que los códigos numéricos también comiencen siempre con la barra "/". Después de la barra, una palabra puede expresar el significado de los siguientes dígitos. Por último, sigue la fecha cuantitativa real. Por ejemplo, si queremos incluir la edad de nuestros interlocutores en las entrevistas, utilizamos la notación "/age: 18" para registrar la información de que la persona tiene 18 años como fecha disponible.

(4) Los *códigos de control interno* del programa o "interruptores" ocupan una posición especial en los archivos. No contienen ninguna información sobre los datos o sus productores. Por lo tanto, no desempeñan ningún papel en el proceso de interpretación. En la versión actual de AQUAD, sólo está definido un código de control para el procesamiento de archivos de texto:

\$no contar

En los archivos de texto, este código de control excluye los segmentos marcados con "\$no contar" del análisis a nivel de palabras individuales (recuento, búsqueda de palabras clave, etc.).

(5) Los *códigos de hablantes* comienzan con los dos caracteres definidos "/"\$ y se comportan como una mezcla de códigos de perfil y códigos de control. Prácticamente dividen un fichero en subfichas para determinados análisis ("búsqueda" y "análisis de tablas"). Los llamamos "códigos de hablantes" porque se introdujeron para apoyar el análisis de las discusiones de grupo. Si las partes de la oración se asignan a los miembros individuales del grupo con códigos de hablante, el texto puede analizarse tanto en su conjunto como en las partes asignadas individualmente (por ejemplo, /\$Juan, /\$María, etc.). También se pueden utilizar estos códigos para distinguir respuestas cortas o declaraciones de muchas personas en un mismo archivo (por ejemplo, la transcripción de las respuestas a algunas preguntas abiertas). En otras palabras, los códigos de hablante asignan secciones de archivos a personas específicas.

Atención: En el análisis, los códigos de los hablantes sólo tienen efecto

- cuando se buscan códigos (es decir, no cuando se cuentan los códigos - esto se maneja para los diferentes hablantes con el análisis de la tabla);
- en el primer nivel de análisis de la tabla, es decir, como definiciones de columna de nivel superior.

(6) Los *códigos de secuencia* se corresponden en principio con (1) los códigos conceptuales, en el sentido de que representan "etiquetas para eventos distinguibles...". La diferencia con los códigos conceptuales simples es que éstos definen un segmento de datos, que a su vez está determinado por dos o más códigos conceptuales en una secuencia definida. Si, por ejemplo, al codificar los segmentos de datos en los que se da un consejo a alguien (véase (1), código conceptual "consejo"), nos encontramos a menudo con que el consejo no es aceptado, sino rechazado, podríamos distinguir dos secuencias típicas de consejo: "consejo - aceptación" y "consejo - rechazo". Podemos entonces insertar el código "Asesoramiento - Aceptación" (o simplemente "Asesoramiento +") en el primer caso, y el código "Asesoramiento - Rechazo" (o "Asesoramiento -") en el segundo. El segmento de datos codificado y vinculado

internamente va entonces desde el inicio del asesoramiento hasta el final de la reacción de aceptación/rechazo de la persona asesorada.

Volvamos al principio del proceso: En la fase de reducción de datos, todos los archivos deben estar codificados. Sin embargo, a continuación hay que tener en cuenta las diferencias al tratar con archivos de texto, sonido, vídeo e imagen. Las reglas básicas de codificación con AQUAD se presentan en detalle utilizando archivos de texto como ejemplo; las características especiales de los otros tipos de archivos se presentan después.

6.2.2 Selección de archivos de texto para la codificación

Descubramos un poco más sobre la codificación. Tras hacer clic en "Análisis CUAL_itivo de contenido", se abrirá un cuadro que muestra los nombres de todos los archivos de texto de su proyecto actual.

Antes de codificar, se tiene que decidir qué archivo de texto se quiere editar. Los nombres de los archivos se muestran en la ventana. Seleccione uno de ellos haciendo doble clic. Para aportar un ejemplo de cómo trabajar con AQUAD, se encuentra en la carpeta principal del programa cuatro archivos: "entre_1.txt", "entre_2.txt", "entre_3.txt" y "entre_4.txt" en el formato txt (ANSI), contruidos de las transcripciones de dos entrevistas con profesores del aula en una publicación de Marcelo (1991) y de unos ejemplos del pensamiento de profesores sobre sus dilemas en una publicación de Zabalza (1991).

Una "X" delante de los nombres de los archivos indica que ya se ha codificado dentro de estos textos, es decir, que ya se ha trabajado con ellos.

Si decides seleccionar "entre_1.txt", AQUAD carga la ventana de codificación con varios botones de función. Más adelante se hablará de esto. Antes de que haga clic en "entre_1.txt" ahora, permítanos recordarle de nuevo que las líneas de texto no deben ser demasiado largas.

6.2.3 Reglas para la inserción de códigos

Cuando haya seleccionado el archivo que desea codificar, se abrirá una ventana con una tabla gris vacía a la derecha. Si ya ha trabajado con este archivo, la tabla muestra las codificaciones anteriores.

La primera columna "M" de la ventana de códigos mostrará posteriormente si se han escrito memos para determinados códigos, y reproducirá estos memos haciendo clic en ellos. Las columnas "desde" y "hasta" indican sobre qué líneas del texto se extiende el segmento de texto codificado.

A la izquierda se ve la ventana de texto resaltada en amarillo claro:

Análisis CUALitativo de contenido - entre_1.txt

M	desde	hasta	Código
X	1	5	Sno contar
	1	5	/SEntrevistador
	6	9	DP exp. doc. previas
	6	20	/SProfesor
X	11	16	CA procedencia
	11	16	clase - alumnos
	11	16	IC ambiente
	18	19	CA conocim. previos
	18	19	CA rendimiento
	18	19	clase - alumnos
	21	22	Sno contar
	21	22	/SEntrevistador
	23	32	/SProfesor
	23	32	CA conocim. previos
	23	32	clase - alumnos
	33	37	Sno contar
	33	37	/SEntrevistador
	38	48	/SProfesor
	38	48	IC infraestructura
	49	50	Sno contar
	49	50	/SEntrevistador
	51	65	/SProfesor

Para codificar cualquier segmento, mueva el puntero del ratón a la izquierda de la ventana de texto hasta el número de las primeras líneas del segmento de texto que desea codificar. Ahora mueva el cursor hasta el número de la última línea de este segmento de texto mientras mantiene pulsado el botón izquierdo del ratón y luego suelte el botón del ratón.

Análisis CUALitativo de contenido - entre_1.txt

Códigos

Buscar

Sustituir

Palabra clave

Buscar

Linea temporal

Ahora suceden dos cosas: En primer lugar, vemos a la izquierda en la ventana de texto que se ha marcado el segmento de texto seleccionado (líneas 1 - 5). En segundo lugar, se abre a la derecha la ventana de introducción de códigos. Vemos en las entradas automáticas que el segmento de texto está en las líneas 1 hasta 5. Desde el registro de códigos podemos ahora seleccionar un código utilizado anteriormente haciendo clic sobre él o, si no se ha introducido todavía ningún código adecuado, escribir un nuevo código en una de las tres líneas de entrada. Concluye correctamente: a cada segmento de texto se le pueden asignar hasta tres códigos a la vez en el mismo paso de trabajo. Al hacer clic en "Guardar", el código o los códigos seleccionados se transferirán a la tabla de códigos.

Por el contrario, si hace clic en uno de los códigos de la columna 4 ("Código:") en la ventana de códigos, el segmento de texto así codificado se marca a la izquierda en la ventana de texto. Al mismo tiempo, este segmento de texto se muestra por separado en una ventana con fondo verde. Allí puede llamar al editor de texto ("notepad")

pulsando el botón derecho del ratón y copiar el segmento seleccionado (para su posterior inserción en otros archivos de texto), imprimirlo, etc.

También puede llamar el registro de código de la derecha para copiarlo o imprimirlo en el editor de texto haciendo clic en cualquier lugar del registro con el botón derecho del ratón.

Ha resultado útil dotar a cada nuevo nombre de código de una breve definición. De este modo, siempre podrá recordar rápidamente el significado o presentar posteriormente su sistema de codificación de forma muy sencilla con el código, la definición y el ejemplo de texto. Para ello, utilice la función "*Memo*"; se activa al hacer clic en la primera columna de esta entrada de código (véase más arriba).

Por favor, recuerde cuando marque algo en el campo de texto: El texto es su base de datos, no puede cambiarlo durante el análisis! AQUAD sólo muestra el texto, pero no permite editarlo. Sin embargo, puedes copiar partes del texto (o todo el texto) y trasladarlo a una nota: Para copiar, basta con pulsar el botón derecho del ratón en la ventana de texto. Todo el texto aparecerá en el editor de texto ("*notepad*"), donde puedes seleccionar las partes críticas del texto, copiarlas (y luego pegarlas en un memo, por ejemplo), imprimir las, etc.

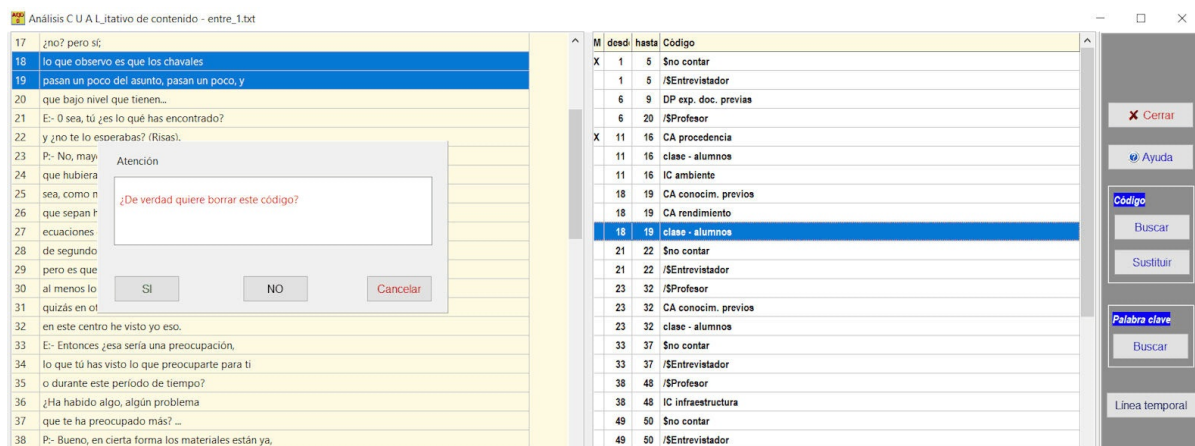
6.2.4 Cómo eliminar o sustituir los códigos

Si pulsa "*Cancelar*" (en lugar de "*Guardar*") en la ventana de codificación, su última entrada será simplemente ignorada. Si descubre un error de escritura antes de "*Guardar*", basta con hacer doble clic en la pequeña ventana de entrada de los nombres de los códigos. El nombre desaparece y puede introducir un nuevo nombre. Por supuesto, puedes moverte en este cuadro con el cursor, borrar o sobrescribir caracteres, como en un programa de texto.

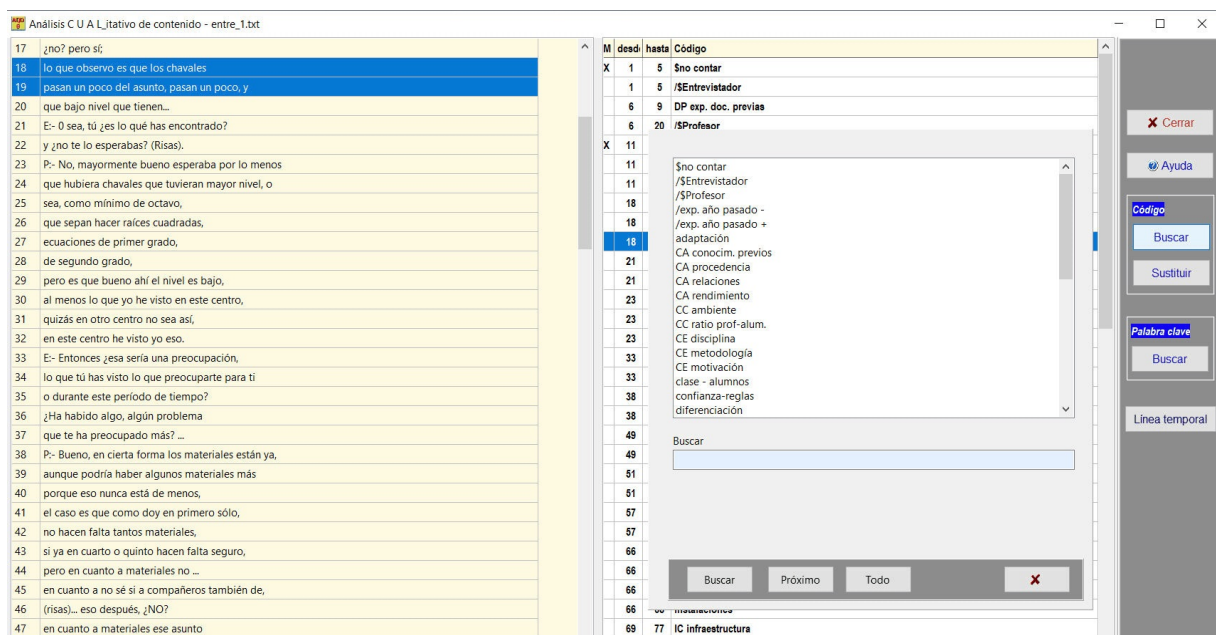
Pero, ¿qué hacer si más tarde -después de haber adjuntado el código al segmento de texto pulsando el botón "*Guardar*"- se descubre que algunos códigos contienen errores tipográficos, son inapropiados o incluso engañosos? Por supuesto, luego querrá eliminarlos. Aquí está la sugerencia sobre cómo proceder. Más adelante, le mostraremos una forma más elegante pero también más peligrosa de buscar y reemplazar automáticamente los códigos.

- Busque el código crítico en la tabla de codificación de la derecha y haga clic delante de él en la columna de la tabla asignada "*desde*" o "*basta*". A la izquierda de la ventana de texto aparecerá una ventana gris (véase la figura siguiente): "¿Desea realmente eliminar este código?"
- Para eliminar el código seleccionado, haga clic en "*SI*".
- A continuación, puede volver a introducir el código corregido.

Si sólo quiere hacer pequeñas correcciones, tiene sentido marcar primero el segmento de texto codificado de nuevo, seleccionar el código incorrecto en la ventana de entrada de códigos, corregirlo en la línea de entrada y guardarlo. A continuación, borre el código defectuoso como se ha descrito. De este modo, no tendrá que volver a escribir todo el código.



Sin embargo, borrar con esta función se convierte en una faena si quieres borrar más de una entrada de código. Por lo tanto, hay un campo "Código" en el lado derecho que contiene dos botones, "Buscar" y "Sustituir". Como ya has adivinado, el botón "Buscar" es útil en este caso:



Aquí tampoco tiene que escribir el código que busca en el campo de entrada, sino que hace clic en el registro de código. Allí se transfiere el criterio de búsqueda a la ventana de entrada "Búsqueda" haciendo clic en ella. Sólo queda iniciar la búsqueda haciendo clic en "Buscar" en la parte inferior.

Esto le llevará de un segmento de texto al siguiente para el que se utilizó este código en particular. Lo mejor es hacer clic en "Todos" para obtener una visión general en una lista de los lugares donde se encontró el código. Con la rutina "Borrar código en cada archivo" (seleccionada en "Tratar códigos" en el menú principal) puede decidir si quiere eliminar todas las instancias del código crítico o si prefiere tomar la decisión caso por caso. Precaución: Esta rutina hace exactamente lo que su nombre indica - es mejor guardar siempre una copia de seguridad de sus códigos en otro directorio.

Puede proceder de forma análoga para sustituir los códigos por otros, ya sea individualmente tras seleccionar "Sustituir" en el campo de botones "Código" o con la rutina "Sustituir código en cada archivo" (seleccionar en "Tratar códigos" en el menú principal).

6.2.5 Cómo obtener una visión general de los códigos

Hay muchas razones por las que es necesaria una visión general más amplia que la que ofrece la pantalla. AQUAD ofrece tres formas de obtener una visión general de los códigos:

Resumen de todos los códigos utilizados en el proyecto.

Como se ha mencionado anteriormente, AQUAD muestra el registro de códigos para ofrecerle una visión general de los códigos utilizados en su análisis en ese momento. Durante la codificación, Aquad mantiene un registro y registra todos los nuevos nombres de código introducidos. Cada nuevo nombre de código se ordena automáticamente por orden alfabético en el registro de códigos. Si quiere volver a utilizar un código, puede seleccionarlo haciendo clic en él.

Si hace clic con el botón derecho del ratón en el registro del código, éste se transfiere al editor y puede copiarse desde allí (para pegarlo en otros textos) o imprimirse.

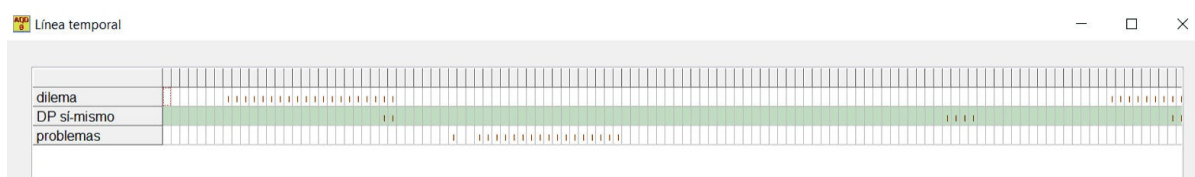
Resumen de los códigos asociados a una línea específica

La columna "Texto" muestra el archivo de texto actual. Los segmentos a los que ya ha adjuntado uno o más códigos se resaltan cuando hace clic en una de las columnas de códigos (véase más arriba). Los números de línea donde comienzan los segmentos de texto aparecen en la columna "desde" de la tabla de códigos. Todos los códigos con el mismo número de inicio de segmento se han asignado al mismo segmento de texto.

Resumen de los códigos de un fichero en contexto secuencial

Una selección de códigos de interés se resume en una lista de códigos (véase en -> "Búsqueda" -> "Catálogo de códigos" -> "Crear un catálogo de códigos") y luego se hace visible dentro de las ventanas de codificación de los proyectos de texto, audio y vídeo tras pulsar el botón "Línea temporal" en su secuencia alterna.

Dado que el propósito de la función es hacer que las secuencias de codificaciones sean claramente manejables, el número de códigos en la pantalla está limitado a 50. Sin embargo, debe tratar de establecer un número menor de códigos críticos si quiere aprovechar de esta función.



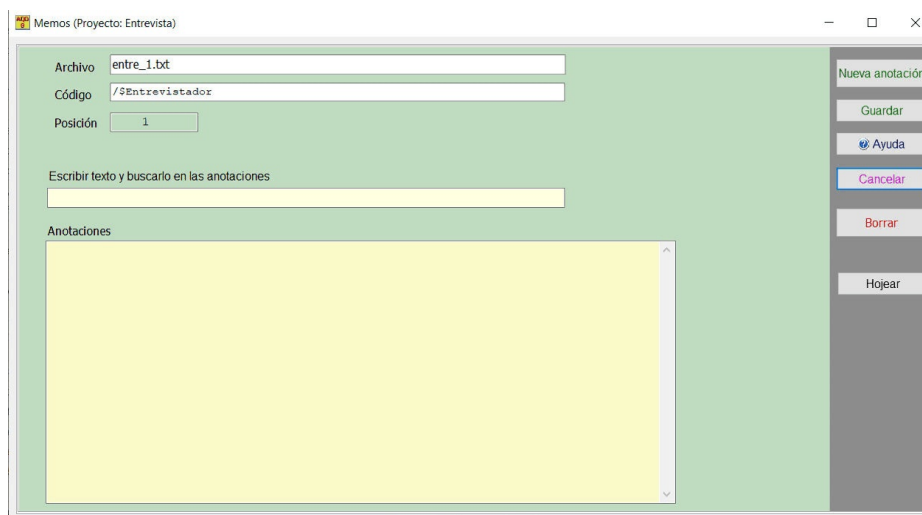
Arriba puede ver el resultado de la selección de códigos del ejemplo entre _2.txt, que se resumió en un catálogo de códigos (*.cco) como se ha descrito anteriormente: Si se hace clic en el nombre de la lista de códigos críticos, AQUAD crea una fila de la tabla para cada código de la lista, en la que su aparición se marca con bloques negros

en relación con el uso de los demás códigos. Las unidades de representación son líneas de texto. En el caso de las grabaciones de audio, las unidades de representación son los segundos, y en el caso de las grabaciones de vídeo, 25 fotogramas cada uno (también equivalente a un segundo). Esto y los errores de redondeo en la conversión de las unidades conducen naturalmente a la superposición de secciones adyacentes de diferentes códigos, pero la posición de los códigos individuales en la secuencia de codificaciones se hace fácilmente comprensible.

6.2.6 Cómo añadir comentarios a los códigos

Si quieres, puedes añadir un comentario a cada código. En teoría, los comentarios pueden ser tan largos como quieras; en la práctica, hay un límite debido a la capacidad de memoria de tu ordenador. Los comentarios se añaden en forma de "memos". Son útiles, por ejemplo, cuando se empieza a codificar con códigos muy generales y ya se prevé que la diferenciación será necesaria a medida que se trabaje con los textos. En una fase posterior, pueden convertir fácilmente estos comentarios en códigos diferenciados. Ahora resulta interesante la primera columna de la ventana de códigos, que lleva por título "M" de "Memo".

Si hace clic directamente en esta columna, se activa la función de notas para la línea en cuestión. Una "X" indica que ya existe al menos una nota para esta codificación que comienza en esta línea. Al hacer clic, se abre la ventana de notas, una ventana vacía si aún no se ha registrado ninguna "X".



Para búsquedas posteriores, y también cuando se selecciona la opción "Memo" del menú principal, se pueden definir tres características para cada memo:

- el nombre del archivo al que pertenece la nota;
- el código que le pertenece;
- la línea de texto a la que se refiere la nota.

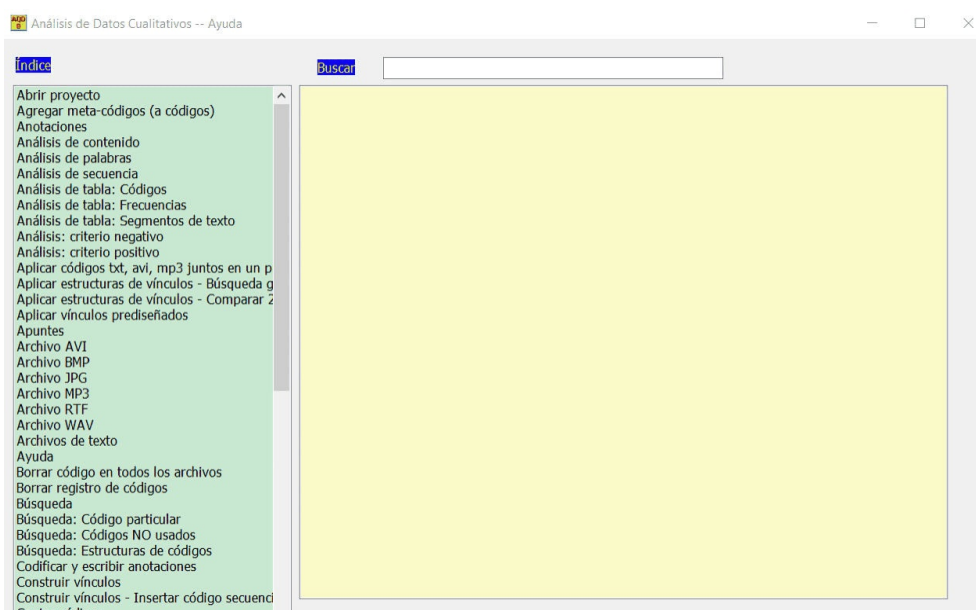
En la ventana bajo el título amarillo *"Escribir texto y buscarlo en las anotaciones"*, puede introducir una palabra clave o una parte coherente del texto, que luego se buscará en todos los memos. Por ejemplo, al codificar, introducimos las definiciones de los códigos como memos, anteponiendo cada vez (por ejemplo) el título *"Definición del código"*. Con esta palabra clave, podemos compilar todas estas definiciones en un instante si seleccionamos la función *"Hojea"* (derecha).

La gran ventana amarillenta no necesita explicación: aquí es donde escribes o importas tu nota. El botón derecho del ratón le permite acceder a copiar, cortar, pegar, etc. Los dos botones de la parte superior derecha hacen lo que indican: puedes utilizarlos para introducir nuevas notas y borrar las existentes.

Los botones de la derecha se explican por sí mismos; sólo "Hojear" puede necesitar una explicación adicional: El desplazamiento en combinación con las funciones – como la búsqueda de un texto específico que se refiera a un código concreto – suele realizarse mediante la función de memo del menú principal. Puede utilizar la función para obtener una visión rápida de las notas con un contenido específico.

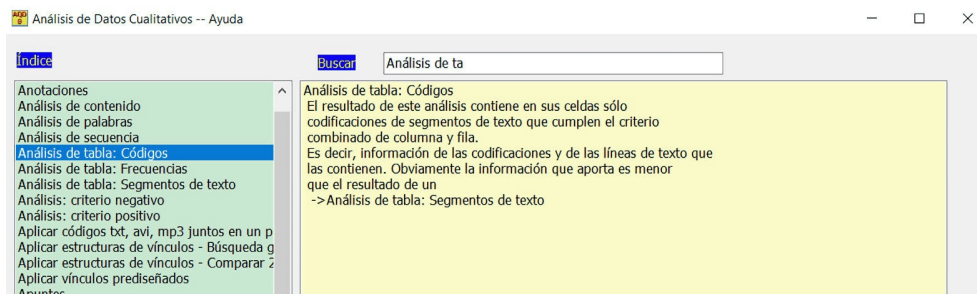
6.2.7 Cuando se necesita un poco de ayuda

No siempre querrá buscar en todo el manual si el significado de una opción del menú no le resulta del todo claro. No hay problema - AQUAD ofrece ayuda general para problemas comunes en el menú principal (ver la última opción del menú principal) y ayuda contextual dentro de las diferentes ventanas. Digamos que quieres ser más específico sobre qué tipos de archivos de texto acepta AQUAD. Hagamos clic en "Ayuda" en el menú principal, luego en "Contenido" en el submenú y veamos qué sucede:



A la izquierda de la figura, una ventana verde muestra la lista de palabras clave disponibles. El texto de ayuda asociado aparecerá a la derecha cuando hagamos clic en una palabra clave. Para obtener la información sobre "Análisis de tabla: Códigos", podemos tirar de la barra de desplazamiento de la izquierda hasta que aparezca la palabra clave deseada, o podemos escribir la palabra clave en la línea de entrada "Buscar" de la parte superior. Observará que mientras escribimos las palabras "Análisis de ta", la lista se desplaza hacia abajo y la palabra que buscamos aparece resaltada. Ahora hacemos clic en la palabra marcada y el texto de ayuda aparece a la derecha (ver página siguiente).

Al pulsar sobre las secciones del texto de ayuda marcadas con flechas, se muestran sugerencias para obtener información adicional. En este ejemplo, dicho enlace está disponible con la palabra clave "-> Análisis de tabla: Segmentos de texto". Lo mejor es llamar a la función "Ayuda" del menú principal y explorar simplemente las posibilidades una vez.



6.2.8 ¿Por qué no dejar que Aquad se encargue de la búsqueda: codificación semiautomática?

Las funciones de búsqueda de AQUAD ofrecen la posibilidad de buscar por ordenador palabras clave en el texto, una especie de codificación semiautomática. Al codificar, hacemos clic en el botón "Buscar" en el campo "Palabra clave" de la derecha, escribimos una palabra clave (o una secuencia de palabras o partes de palabras) en la ventana que aparece, y luego le decimos a AQUAD que recorra el texto.

Así, en nuestro texto de ejemplo, podría interesarnos dónde se dice algo sobre los alumnos. Podemos introducir la palabra clave "alumnos" y hacer clic en "Buscar". Como esta función de búsqueda no es sensible a las mayúsculas y minúsculas, obtendremos como resultado en el texto entre_2.txt todas las apariciones de "alumnos" si hemos elegido la primera de las dos opciones adicionales. Hay dos ajustes

- palabras completas
- partes de palabras

Si elegimos la segunda opción e introducimos de solamente "alumn" como criterio, podemos encontrar no sólo "alumnos" en plural, sino también en singular o también las formas gramaticales femeninas. Un procedimiento tan sencillo suele ayudar a encontrar rápidamente segmentos de texto críticos.

Además, se aplica lo siguiente: *¡siempre debemos comprobar e interpretar todos los pasajes encontrados!*

¿Qué podemos hacer si varias palabras clave son relevantes, o si varias palabras clave definen un campo de significado que es importante para el análisis? Supongamos que queremos averiguar en un texto que los hablantes expresan percepciones más que pensamientos e ideas. Podemos confeccionar una lista de palabras o tallos de palabras relevantes como oír, oído, ver, ojo, etc. Desgraciadamente, la opción de palabras clave en esta parte del programa nos obliga a introducir una palabra clave tras otra y a que se compruebe.

Es más rápido encontrar todas las palabras juntas utilizando la función de palabras clave del módulo "Búsqueda". Si quiere tener una visión rápida de las palabras clave en todos sus archivos (sin cargar uno tras otro), puede utilizar la opción "Contar palabras" en el mismo módulo (más sobre esto en la sección 7.5 dentro de este capítulo y en el capítulo 9).

6.2.9 Si desea imprimir archivos de código

La lista de códigos que aparece a la derecha de la ventana de *Análisis CUAL_itivo* siempre nos ofrece una visión general de nuestros códigos, pero sólo en la pantalla cuando tenemos AQUAD abierto. ¿Cómo se procede si se desea una impresión de las codificaciones?

- Active la función "*Análisis CUAL_itivo de contenido*";
- seleccione el archivo de texto para el que desea imprimir los códigos;
- haga clic en cualquier lugar de la columna "Código" - con el **botón derecho** del ratón. Toda la lista se carga inmediatamente en el editor ("*notepad* + +") y se muestra allí.
- Allí, seleccione la opción "*Imprimir*" en el grupo de menú "*Archivo*".

6.2.10 ¿Qué es el registro de códigos y por qué querría eliminarlo?

En primer lugar, algunas distinciones terminológicas importantes:

- Las codificaciones están compuestas por el número de texto, el número de la línea de inicio y final de un segmento de texto, el código real.
- Los códigos se almacenan en archivos de códigos (*.cod).
- Cuando hablamos de un registro de códigos, nos referimos a la lista alfabética de códigos utilizados en todos los archivos de un proyecto.

En varias ocasiones, AQUAD quiere saber a cuál de los muchos códigos quieres referirte para el análisis. Esto ocurre cada vez que se adjuntan códigos a los archivos de texto. Si prefiere seleccionar los nombres de los códigos haciendo clic en ellos en lugar de escribirlos cada vez, AQUAD necesita mantener un registro de los códigos utilizados. Otro caso es cuando se quiere tener una visión general de qué códigos se utilizaron para reducir qué textos y con qué frecuencia. Como requisito previo a dicho análisis de frecuencia, AQUAD quiere saber primero qué códigos debe buscar y contar: ¿todos o sólo algunos críticos para tu pregunta? Incluso si se quiere contar todo, AQUAD necesita saber qué códigos se han utilizado; de lo contrario, no será posible determinar si algunos de los códigos pueden no haberse utilizado nunca en algunos de los textos.

Así que necesitas una lista de todos los nombres en clave. Por lo tanto, mientras codificas, AQUAD hace una especie de contabilidad de los nombres de código que introduces mientras trabajas. Cada vez que se utiliza un nuevo nombre de código, éste se introduce automáticamente en un registro de códigos ordenado alfabéticamente. Si sólo unos pocos códigos son significativos desde el punto de vista del tipo de análisis, puede simplemente seleccionar los códigos de esta lista de códigos creada por AQUAD haciendo clic en ellos. Si casi todos los códigos son relevantes para el análisis, como en el caso del recuento, elimine sólo los pocos códigos de la lista que no desee utilizar. Por supuesto, su registro de código permanece intacto con este procedimiento; usted siempre trabaja con una copia del directorio original.

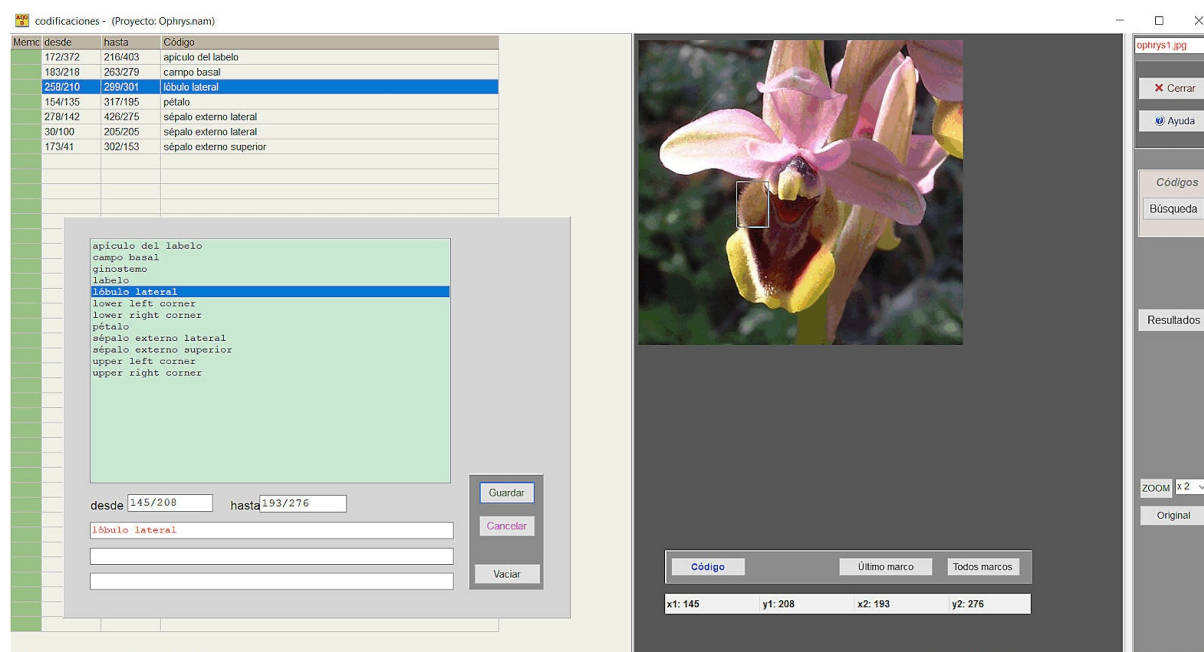
AQUAD crea la lista automáticamente y también la mantiene actualizada de forma automática, así que ¿por qué ibas a querer eliminarla? Pues bien, podría surgir la siguiente dificultad: Si se trabaja con la posibilidad de sustituir los códigos por metacódigos – y se debe jugar con esta posibilidad – puede ocurrir que después algunas entradas de código aparezcan dos o más veces en el registro de códigos. Algo similar puede ocurrir si modificas los archivos de código con un editor de texto fuera de AQUAD.

Esto no causa ningún problema en términos de funciones, pero contradice la lógica y la estética de un registro de códigos. Para limpiar el registro de códigos, basta con borrarlo (opción del submenú *Editar Códigos*) y luego restaurarlo (última opción del mismo submenú). El nuevo registro de códigos se crea a partir de todos los códigos asociados a los archivos que ahora lee AQUAD para su control.

6.2.11 Codificación de archivos de imagen

Las ventanas y funciones del programa para codificar archivos de imagen son las mismas que las descritas anteriormente para codificar textos.

En la siguiente captura de pantalla vemos la lista de codificación a la izquierda, la imagen a codificar a la derecha. Los segmentos del archivo se marcan con el ratón (pulse el botón izquierdo del ratón en la esquina superior izquierda o inferior derecha del área a marcar, manténgalo pulsado, arrastre el marco, suelte el botón): En la foto de una flor de *Ophrys*, se acaban de marcar los pétalos (véase la ilustración); AQUAD introduce las coordenadas de la imagen y, tras hacer clic en el botón "Código", abre la ventana de codificación en la que se puede seleccionar el código adecuado del registro o introducirlo de nuevo. Haga clic en "Guardar" para guardar toda la codificación (coordenadas de la sección de imagen y código) en el archivo de código.

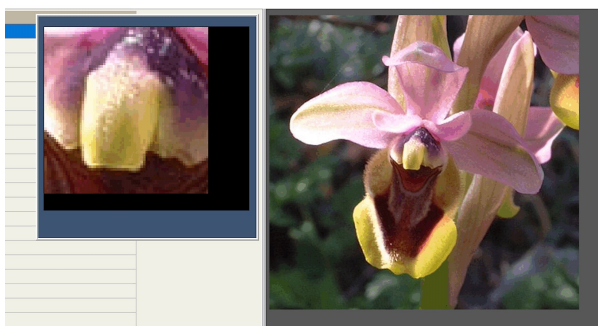


A la inversa, si hace clic en un nombre de código en la lista de codificación, el segmento de imagen correspondiente se resalta a la derecha. También es posible borrar las entradas de código tras hacer clic en las coordenadas de un código, como cuando se trabaja con textos.

Sin embargo, a diferencia de los otros tipos de archivos, no es posible buscar estructuras de código específicas y enlaces de códigos en los códigos de los archivos de imagen. El motivo es la localización algo complicada de los segmentos de datos mediante la especificación de las coordenadas x/y de los puntos de la esquina superior izquierda e inferior derecha. Hasta ahora, los usuarios no han expresado la necesidad de estas funciones al analizar los datos de

las imágenes. Sin embargo, las estructuras de codificación y las hipótesis de vinculación de los textos que están disponibles como archivos de imagen (por ejemplo, los textos manuscritos escaneados) pueden comprobarse con un pequeño truco (véase la sección siguiente).

Dado que todas las imágenes en AQUAD se reducen en tamaño al formato de ventana si son más grandes en el original, se dispone de una función de lupa para ver los detalles (ver más abajo). Moviendo el puntero del ratón sobre la imagen, se selecciona el área que se desea ampliar. El botón "Zoom" sirve para encender y al segundo pulso para apagar la función de lupa.

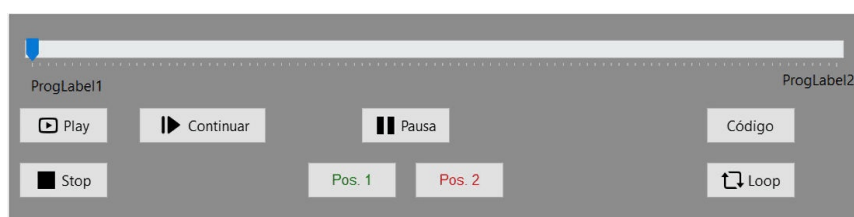


Una función poco utilizada que determinaba información cuantitativa sobre los detalles de la imagen se eliminó en la versión 8, pero puede volver a instalarse rápidamente si es necesario: Para algunas preguntas de investigación, la información cuantitativa de las imágenes puede ser interesante: Por ejemplo, ¿qué tamaño tenía el padre en los dibujos de los niños en comparación con la madre y los hermanos? ¿Qué proporción de la superficie ocupa una determinada área de la imagen en los dibujos de diferentes niños? Para responder a estas y otras preguntas similares, en Aquad 7 (con la función "Calcular altura/anchura" en el grupo de funciones "Herramientas" del menú principal en "Códigos de imagen") se podía calcular la altura y la anchura de los segmentos de imagen seleccionados (selección basada en los códigos) como la diferencia de las coordenadas de píxeles verticales u horizontales e insertarlas en los códigos.

6.2.12 Codificación de archivos de audio

Trabajar con el reproductor multimedia

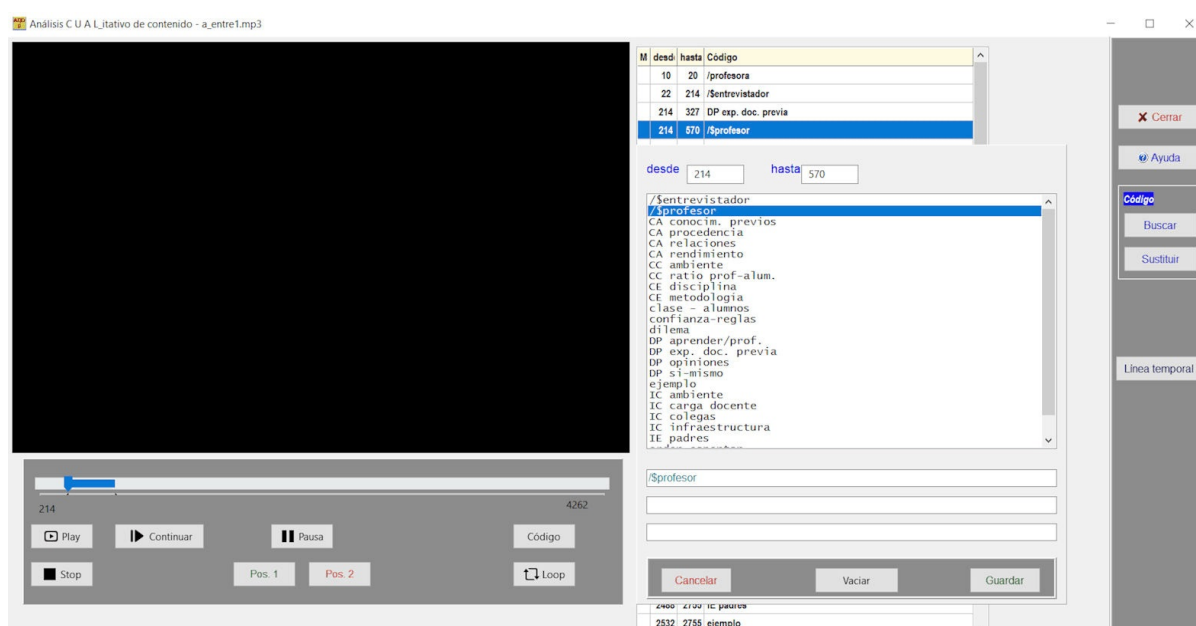
Las ventanas y funciones del programa para la codificación de grabaciones de audio se corresponden con lo descrito anteriormente para la codificación de textos. Sin embargo, al codificar, debe utilizar el control del reproductor multimedia para desplazarse por los datos.



En la figura siguiente, vemos las columnas de los códigos a la derecha: "Memo" para los memos ("X" señala que los memos se han añadido; aquí, hasta ahora no hay apuntes), "de" - "a" y, finalmente, los propios códigos. La

unidad de grabación de sonido es la décima de segundo; internamente, el reproductor multimedia trabaja con milisegundos como unidad, por lo que puede haber pequeñas desviaciones entre el marcaje y el inicio/final de la secuencia de grabación, pero no son molestas en las grabaciones habituales de las entrevistas. La captura de pantalla muestra el código de altavoz marcado "/\$profesor", el segmento de datos asignado va de "214" a "570".

Cuando el profesor comenzó a hablar en este pasaje, se pulsaron los botones "Pausa", luego "Pos1" (-> posición 1) para marcar el comienzo (214), luego "Siguiente" y al final de esta sección de la entrevista, cuando el profesor dejó de hablar (570), los botones "Pausa" y luego "Pos 2". La posición actual respectiva del reproductor en el archivo se muestra a la izquierda bajo el deslizador y se transfiere a la ventana de codificación - aquí vemos la posición de parada "570" después de hacer clic en la posición (final) 2.



La introducción de códigos simples o múltiples para el segmento de datos marcado corresponde exactamente al procedimiento cuando se trabaja con datos de texto.

Los segmentos marcados entre las posiciones 1 y 2 (véase más arriba) pueden reproducirse una y otra vez con el botón "Loop". El ajuste es posible directamente cambiando los números de los campos "desde" y "hasta".

Haciendo clic en un código de la columna correspondiente de la ventana de codificación, puede reproducir directamente la sección codificada.

Los movimientos rápidos a través de la grabación se realizan mejor con el control deslizable de la parte superior del panel de control. Haga clic en la lengüeta del deslizador, mantenga pulsado el botón del ratón y mueva el deslizador a la izquierda o a la derecha con el ratón. El valor del contador cambia análogamente. Los botones "Pos1" y "Pos2" mantienen el principio y el final, respectivamente, de un segmento de datos a codificar; el valor del contador se transfiere a la ventana de codificación. Al hacer clic en una posición (columnas "desde" y "hasta") en la tabla de codificación, se abre una pequeña ventana a la izquierda en la que puede volver a eliminar la entrada (con la respuesta "Sí"). A diferencia de lo que ocurre cuando se trabaja con textos, en esta ventana se puede modificar manualmente la posición inicial y/o final del segmento codificado en las dos líneas de entrada "desde" y "hasta".

Cuando se buscan códigos (véase el trabajo con archivos de texto), la ubicación donde se encontró el código se muestra en una ventana verde a la izquierda en cada caso en la lista de codificación. A continuación, cierre la ventana de búsqueda y haga clic en el código que aparece a la izquierda en la ventana de codificación de la derecha: el pasaje codificado se reproduce inmediatamente.

Lo mejor es abrir el proyecto de muestra "a_Interview1" (grabación de audio de las transcripciones del proyecto "Interview") y jugar con las posibilidades.

Estrategia general de codificación de archivos de audio

A diferencia de los textos transcritos, no se puede hojear rápidamente el contenido de los archivos de audio. Por lo tanto, mientras no conozca bien el contenido, necesitará mucho tiempo si quiere encontrar ciertos pasajes. Por lo tanto, recomendamos una estrategia de codificación que ha dado buenos resultados en nuestro propio trabajo:

En primer lugar, estructurar el archivo formalmente, por ejemplo, según los cambios de orador. En el proyecto "a-entrevista1.mp3", los segmentos de archivo de los profesores entrevistados y del entrevistador se marcaron primero con códigos de hablante (/ \$profesor, / \$entrevistador).

Dentro de los segmentos de los oradores, los énfasis temáticos deben registrarse ya durante esta estructuración formal, ya sea en memos o en códigos preliminares.

Por lo tanto, es aconsejable asignar inmediatamente la codificación de hablante adecuada al principio de un nuevo segmento de archivo, incluso si el límite final del segmento todavía se desconoce en ese momento. De lo contrario, si se utiliza una codificación de contenido temático dentro de este segmento, el límite inicial cambiará - y entonces tendrá que buscar laboriosamente el comienzo del segmento de nuevo. Así que proceda así:

- Al comienzo de un nuevo segmento (cambio de orador, nueva pregunta, etc.), pulse inmediatamente el botón de parada y, a continuación, ajuste "Pos 1" (límite inicial) y "Pos 2" (límite provisional, falso límite final). Haga clic en "Código" y establezca un código adecuado (de hablante).
- Si a continuación conoce la marca final exacta del final del segmento, es decir, en el siguiente cambio de orador o de tema (pulse la tecla de parada, lea el número), haga clic en el código actual (orador) con el límite final incorrecto en la tabla de codificación. Ahora puede transferir el valor numérico en el campo de entrada "basta" borrando y escribiendo o copiando y pegando desde el panel operativo de abajo (marque allí, haga clic en el botón derecho del ratón, seleccione "Copiar").

A partir de los comentarios de un usuario de AQUAD, los siguientes consejos pueden ser útiles: Si al pulsar los botones "desde" o "basta" o "Pausa" o "Stop" en el reproductor de audio no se interrumpe inmediatamente la reproducción – se han reportado tiempos de retardo del orden de segundos –, no busques la causa en AQUAD, sino en la configuración del sistema de tu ordenador. La causa más probable, especialmente en los ordenadores con "sonido a bordo", es que el software del controlador suministrado o instalado por el distribuidor y el hardware no coinciden. La solución más sencilla es consultar el manual del ordenador para averiguar el fabricante del componente de "sonido a bordo" y, a continuación, descargar el último controlador de la página web del fabricante en Internet.

Transcripción de segmentos de archivos importantes

La codificación directa de las grabaciones de audio sin transcripción ahorra mucho tiempo, pero los pasajes importantes no pueden copiarse simplemente en el informe de investigación. Por ello, es aconsejable transcribir los pasajes críticos o significativos, indicando el recuento en el reproductor multimedia, y recogerlos en un archivo de texto adicional (por ejemplo, el sencillo editor notepad.exe).

6.2.13 Codificación de archivos de vídeo

Cuando se trabaja con grabaciones de vídeo, todo lo descrito anteriormente para las grabaciones de audio se aplica de forma análoga. La diferencia decisiva es que ahora encontrará una pantalla de vídeo sobre el panel de mando en la que se reproduce la grabación.

La recomendación para la codificación de datos de audio de realizar primero una pasada de codificación formal (ver arriba) para estructurar la grabación según los cambios de hablante y/o escena se aplica de forma análoga a los datos de vídeo.

6.3 ¿Cómo se editan los textos y las codificaciones?

6.3.1 Cómo editar textos

Esto se aclara rápidamente: ¡en AQUAD no se editan textos en absoluto! AQUAD está concebido como un programa de análisis de textos. Los textos, que constituyen la base de los análisis, no pueden ser modificados durante el proceso de análisis.

Sin embargo, puede haber razones para no ser tan ortodoxo con la base de datos. Es posible que uno descubra uno o dos errores de escritura que prefiere mejorar antes que documentar hasta la publicación de los resultados de la investigación. A veces se encuentra que la información adicional sería útil para la interpretación, como las referencias entre corchetes al comportamiento no verbal en las transcripciones de las grabaciones de vídeo de una discusión. (Sin embargo, con AQUAD 8 de vídeos, es probable que sólo transcriba escenas seleccionadas, pero en general analice el vídeo directamente). Por último, en los análisis a nivel de palabras, es posible que desee excluir determinados pasajes de texto, como las preguntas del entrevistador o la información que identifica a los interlocutores, la hora, el lugar, etc. En todos estos casos, hay que editar la base textual del análisis.

Por supuesto, esto siempre es posible fuera de AQUAD antes de empezar a codificar. Al hacerlo, siga las directrices de preparación de textos para el análisis con AQUAD y no olvide volver a guardar sus textos en formato *txt* después de editarlos con su procesador de textos.

Sin embargo, si ya ha codificado los textos, tendrá problemas después de editarlos. La edición puede haber cambiado la longitud del texto y las posiciones iniciales de los segmentos de texto en las codificaciones ya no coinciden con el texto modificado. A continuación, las codificaciones se colocan incorrectamente en el texto a partir del primer cambio. Por lo tanto, la preparación de los textos para el análisis de contenido debe considerarse y planificarse a fondo antes de trabajar con AQUAD, es decir, antes de la codificación.

6.3.2 Cómo editar los códigos

Recuerde que las codificaciones se componen de los números de línea (línea inicial y final) del segmento de texto codificado y de las palabras mismas del código. Es posible que haya cometido un error al codificar, o que quiera diferenciar un código muy general como "emoción" después e introducir en su lugar "alegría" o "miedo", etc. Estos códigos pueden editarse cuando se codifican archivos de texto, así como cuando se codifican archivos de audio y vídeo o archivos de imagen:

- Dentro de la codificación de "texto", así como de "grafic", "audio" o "vídeo", puede eliminar códigos y añadir otros nuevos.

- Dentro de la codificación de "texto" así como de "grafic", "audio" o "vídeo" hay un botón especial "Buscar" en el campo "Código" en la parte derecha de la ventana. Este botón da acceso a las funciones de búsqueda y reemplazo que le ayudan a editar códigos específicos en su archivo de una vez.

Recuerde: AQUAD sólo permite editar los códigos asociados a un segmento de datos, no permite editar el contenido del segmento!

Recuerde: Si desea comprobar una entrada de código en la lista de codificación que aparece en la pantalla al codificar "texto", "grafic", "audio" o "vídeo", haga clic en el nombre del código para abrir una ventana que le permita confirmar o eliminar la codificación o cambiar los límites del segmento del código al codificar "grafic", "audio" o "vídeo".

Cuando trabajes con archivos de audio o vídeo, haz clic en el botón "Loop" situado en la parte inferior del panel de control del reproductor de audio o vídeo. Esto le permite escuchar o ver el segmento de datos codificado una y otra vez. Cuando se trabaja con archivos de imagen, al hacer clic en el código de la tabla de codificación se marca el segmento de imagen codificado con un rectángulo..

6.4 ¿Para qué sirven los metacódigos?

Si se sigue la estrategia de generalización de la codificación, en la primera pasada por los datos de un proyecto se inventa un gran número de códigos diferentes para captar adecuadamente la variedad de contenidos relevantes para la pregunta de investigación. Sin embargo, pronto se llega a un punto en el que es necesario buscar puntos comunes dentro de la multitud de codificaciones para luego proceder a resumir y estructurar los códigos. Pero también puede enfrentarse a esta necesidad si procede de forma poco estratégica: cuando revisa los archivos por primera vez (los lee, los escucha, los mira), suele encontrar muchos segmentos interesantes que pueden marcarse con muchos códigos significativos. Pronto se corre el riesgo de perder la visión de conjunto y, por tanto, de utilizar códigos diferentes para unidades de significado similares. No se puede evitar poner orden en la larga lista de códigos diferentes. Al hacerlo, se tropieza rápidamente con las posibilidades de edición manual, porque ahora se trata de reorganizar todo el sistema de códigos.

El principio de "comparación constante" contiene un enfoque para resolver estos problemas también. El procedimiento consiste en comparar los códigos para ver cuáles de ellos representan significados similares y, por tanto, pueden entenderse como un grupo de subconceptos que pueden asignarse a un concepto superior más abstracto. Estos códigos de nivel superior se denominan metacódigos en AQUAD. Si cree que los metacódigos podrían proporcionar una visión general y estructurada en su análisis de texto, le recomendamos que imprima el registro de

códigos, es decir, la lista de todos los códigos utilizados hasta el momento (¡haga clic con el botón derecho del ratón en la ventana de registro de códigos!) y marque con diferentes colores todos los códigos que pertenecen juntos según su significado. A continuación, debe buscar el nombre de un concepto superior para cada uno de los grupos de códigos, es decir, el metacódigo.

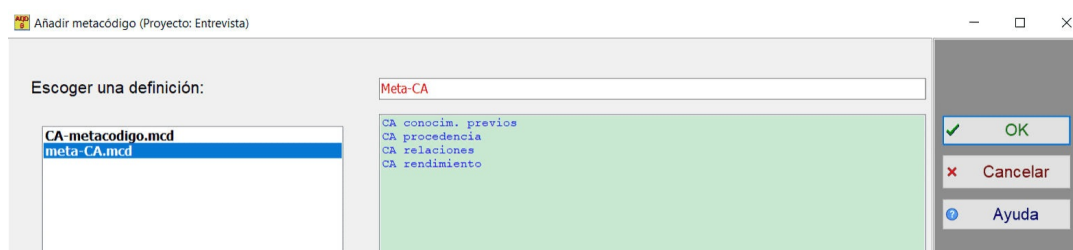
Tras este trabajo preliminar, puede empezar a definir los nuevos metacódigos que se introducirán uno tras otro. Para ello, seleccione la opción "Metacódigos" en el submenú "Editar códigos" y allí "Crear/modificar definición". Al crear una nueva definición, naturalmente se responde a la pregunta de si se quiere aplicar una definición ya disponible con "NO". Ahora aparece una nueva ventana: introduzca primero el nombre del nuevo metacódigo en la línea "Introducir metacódigo". A la izquierda verá el registro de códigos desde el que puede seleccionar todos los códigos que se subsumirán en un nuevo metacódigo haciendo clic en el nombre del código o códigos en el cuadro verde claro (registro de códigos) de la izquierda. Puede transferir los códigos del registro de códigos a la definición de su metacódigo haciendo clic sobre ellos, o devolverlos al registro de códigos en sentido contrario haciendo clic sobre ellos.

Cuando haga clic en "Aceptar", la definición de su nuevo metacódigo, es decir, el nombre del metacódigo y la lista de códigos subordinados, se guardará para su uso posterior. Por lo tanto, posteriormente se le pedirá que escriba un nombre de archivo (sin extensión) bajo el cual se almacenará la definición del metacódigo. Ha resultado útil guardar estas listas de definiciones bajo el nombre del respectivo metacódigo definido con una extensión precedente, por ejemplo, la definición de un metacódigo "asesoramiento" (en el que deben combinarse los códigos "consejo", "asesorar", "ayudar", "sugerir", "proponer", "señalar") como "M_Asesoramiento".

6.4.1 Modificar archivos: Añadir metacódigos

Si en el submenú "Tratar códigos" se selecciona la opción "Añadir metacódigo", el nuevo metacódigo se introducirá en el fichero además de cada uno de los antiguos códigos con el mismo significado. De este modo, se conservan los códigos originales. Sin embargo, su archivo de código se hará cada vez más largo, especialmente si quiere jugar con la posibilidad de generalización introduciendo metacódigos - lo que definitivamente debería intentar. Si utiliza esta opción con frecuencia, es mejor elegir la función "Sustituir códigos por metacódigos", que, sin embargo, garantiza la conservación de sus archivos de código originales para posteriores análisis.

Para añadir metacódigos, primero hay que seleccionar una definición de metacódigo (véase más arriba) de la lista de definiciones guardadas (a la izquierda en la figura), por ejemplo "Meta-CA.mcd", en la que se resumieron todos los códigos que se refieren a los segmentos de la entrevista en los que los profesores hablaron de su clase y sus alumnos en general (CA ...). Al lado, se muestra el nuevo nombre del (meta)código en la parte superior para su comprobación, así como los códigos subordinados al nuevo metacódigo. Si hacemos clic en "OK", este metacódigo se añadirá a los códigos de los segmentos críticos del archivo:



6.4.2 Modificar archivos: Sustituir los códigos por metacódigos

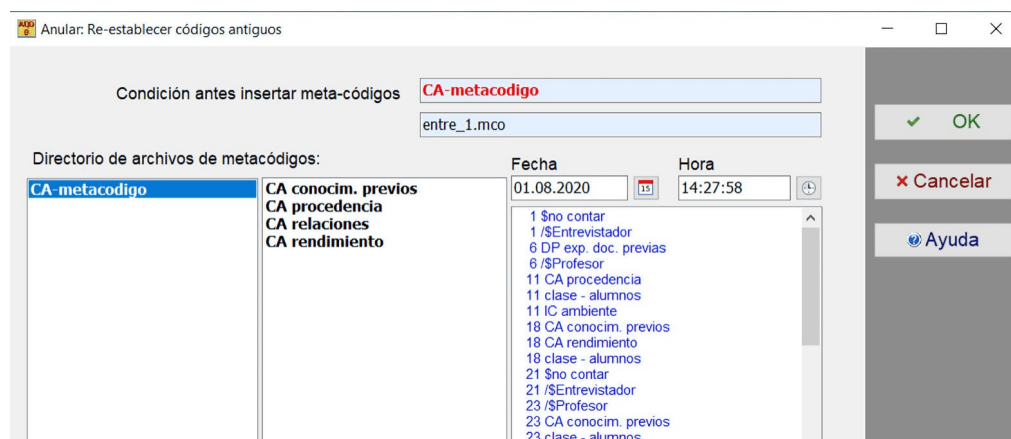
El procedimiento es el mismo que para añadir metacódigos. Sin embargo, tras pulsar "OK", los códigos subordinados desaparecen de la tabla de codificación y son sustituidos por el nuevo metacódigo. Sin embargo, los archivos de codificación antiguos se siguen guardando en el directorio ...\\mco\\ con la fecha y la hora del cambio. De este modo, sus códigos originales permanecen intactos y pueden reactivarse para análisis posteriores si es necesario.

Sin embargo, AQUAD sólo almacena hasta 100 variantes de codificación. A partir del 101. cambio por sustitución, los archivos antiguos – empezando por el más antiguo – se sobrescriben en secuencia. AQUAD te llamará la atención para que puedas guardar manualmente las antiguas codificaciones en otro directorio.

6.4.3 Restaurar los códigos anteriores

Si quiere volver a trabajar con un estado de codificación anterior, seleccione la opción "Metacódigos" en el submenú "Tratar códigos" y allí la opción "Anular: Re-establecer códigos antiguos".

En la ventana que se abre, seleccione una definición de metacódigo cuyos componentes se muestran en la ventana central. Ahora puede restaurar el estado anterior a la aplicación de este metacódigo. En la ventana de la derecha se ve el estado de la codificación del primer archivo de su proyecto, encima la fecha y la hora de la modificación con el metacódigo seleccionado. Al hacer clic en "OK" se cargan los archivos de código válidos hasta ese momento para todos los archivos del proyecto:



Cuando se trabaja con metacódigos, es conveniente utilizar la función de memos/apuntes y llevar un registro en un diario de investigación de cuándo exactamente se introdujeron y utilizaron los metacódigos. Con estos registros, no debería ser difícil recuperar los archivos de códigos correctos entre el máximo de 100 estados de codificación almacenados.

6.5 ¿Cómo insertar varios códigos?

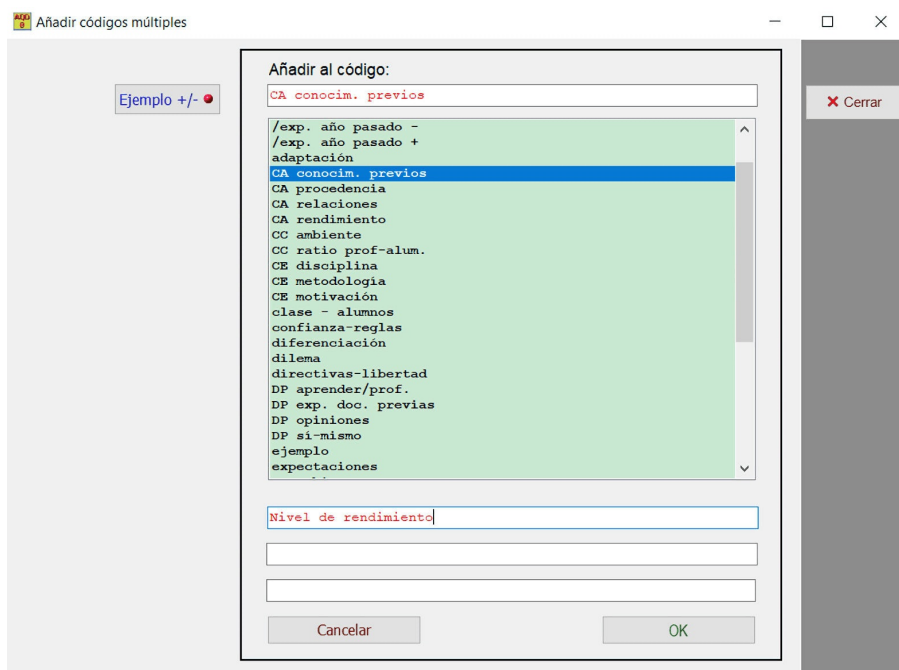
Durante la codificación, se puede marcar un segmento de archivo seleccionado (segmento de texto, secuencia de vídeo, etc.) con hasta tres codificaciones a la vez. Esto es útil, por ejemplo, si quiere marcar los segmentos de texto de un

determinado orador con su código específico de orador (/\$....) y también con un código conceptual y, además, excluirlos del recuento (\$no contar).

Sin embargo, en ocasiones, sólo en el transcurso de la codificación surge la necesidad de codificar adicionalmente algunos segmentos del archivo ya codificados, es decir, de insertar múltiples códigos en los puntos ya codificados. Los segmentos de archivos codificados podrían encontrarse rápidamente utilizando la función "Buscar códigos", pero entonces habría que reiniciar la función de introducción de códigos individualmente *en cada lugar*. Por supuesto, también se podría utilizar la función "Añadir metacódigos" para este propósito de una manera ligeramente modificada y entonces tener el efecto de la inserción automática de código en posiciones previamente definidas.

La situación es más clara con la función especial "Intercalar códigos múltiples" (que se encuentra en el grupo de funciones "Tratar códigos"), porque hace exactamente lo que promete: primero seleccionamos del registro de códigos haciendo doble clic en el código a cuyos segmentos de archivo (por ejemplo, secciones de texto) hay que añadir más códigos múltiples. A continuación, seleccionamos el/los código/s adicional/es o escribimos otros nuevos en las columnas de entrada de abajo si queremos añadir códigos que aún no se han utilizado (véase la ilustración siguiente). Al hacer clic en el botón "Intercalar códigos múltiples" la función comienza su trabajo en todos los archivos del proyecto activo o directorio de archivos.

En la captura de pantalla vemos un ejemplo del proyecto "Entrevista" en el que, por la razón que sea, se decidió durante la codificación asignar posteriormente otros códigos a todos los segmentos de texto que ya habían sido interpretados previamente con el código conceptual "CA conocim. previos".



Si pulsa el botón "Ejemplo +/-" (arriba a la izquierda, letra azul), el programa introduce un código de ejemplo, que se complementa con un código adicional introducido a continuación para cada ocurrencia. De este modo, se podría por ejemplo sustituir posteriormente el código "/\$Entrevistador" por el código de control "\$not contar". Por supuesto, también se puede insertar un segundo o tercer código a la vez. Al pulsar el botón por segunda vez, las entradas de muestra se borran de nuevo y puedes insertar tus propios códigos.

6.6 ¿Cómo se pueden utilizar las palabras clave?

Esto ya se ha explicado anteriormente. Recuerde el botón "*Buscar*" en el campo "*Palabra clave*" de la parte derecha de la ventana cuando codifique textos. Mientras lee sus archivos de texto, puede utilizarlo para activar la búsqueda de palabras clave.

Se hace clic en "*Buscar*" y aparece una ventana especial en la que se introduce la palabra crítica.

La búsqueda se inicia pulsando el botón "*Buscar*" en la parte inferior de la ventana que aparece y saltando de ocurrencia en ocurrencia con "*Próxima*". O bien, si hace clic en "*Todo*" desde el principio, se mostrarán todos los resultados de una búsqueda ejecutada a la vez. Suponiendo que proceda paso a paso ("*Buscar*" y repetidamente "*Próxima*"): En cuanto se encuentra la palabra buscada, la línea correspondiente en la ventana de texto se resalta en color. ¿Quizás le sorprenda ocasionalmente el resultado? Es importante recordar algunas reglas:

- La búsqueda no distingue entre mayúsculas y minúsculas cuando busca palabras clave en los archivos de texto.
- La función de búsqueda tampoco puede volver a juntar palabras separadas. Si desea utilizar la búsqueda de palabras de forma sistemática, deberá desactivar la separación automática al transcribir los textos.
- Si ha marcado la opción "Partes de palabras" en la ventana de búsqueda, la función informa de la cadena introducida independientemente de su posición al principio, dentro, al final de una palabra o como palabra independiente. Por ejemplo, si deja que el programa busque la palabra "para", también encontrará la palabra "aparato", "esparadrapo", "preparar", etc. Un consejo: supongamos que quiere buscar todos los pasajes de texto que mencionen "ventaja" y "desventaja". Si ahora simplemente introduce el componente común "venta" como palabra de búsqueda, AQUAD encontrará lo que busca en una sola pasada, si las palabras aparecen en el texto. Sin embargo, también aparecerían compuestos con "venta" como "ventana" o "ventada".
- Si acepta la opción por defecto "Palabras completas" en la ventana de búsqueda, la función sólo encontrará, por supuesto, las palabras completas que correspondan a su entrada de búsqueda, pero independientemente de su posición en la frase.

Capítulo 7: Cómo trabajar con códigos

7.1 Cómo volver a encontrar los códigos

7.1.1 Requisito previo: crear un catálogo de códigos

Como requisito previo a la búsqueda de códigos en todos los archivos (no sólo en el archivo actualmente abierto) de un proyecto, los códigos a buscar deben estar agrupados en un catálogo. Si hace clic en la opción "*Catálogo de códigos*" del módulo "*Búsqueda*", se abre una ventana que muestra el registro de códigos del proyecto con fondo verde

a la izquierda, y una ventana vacía para el nuevo catálogo de códigos a la derecha. Encima hay una ventana de entrada en la que se introduce un nombre para guardar el nuevo catálogo.

Para seleccionar los códigos críticos como contenido del catálogo, basta con hacer clic en el nombre del código a la izquierda, que se mostrará como contenido del catálogo a la derecha. Por supuesto, puede añadir más de un nombre de código al catálogo. Si selecciona uno por error, simplemente haga clic en el nombre del código no deseado en la ventana del catálogo de la derecha. Se borrará del catálogo y se reinsertará en el registro de la izquierda. No te preocupes: el registro del código en sí no cambia, AQUAD sólo utiliza una copia del registro aquí.

7.1.2 Cómo encontrar segmentos de texto codificados y codificaciones específicas

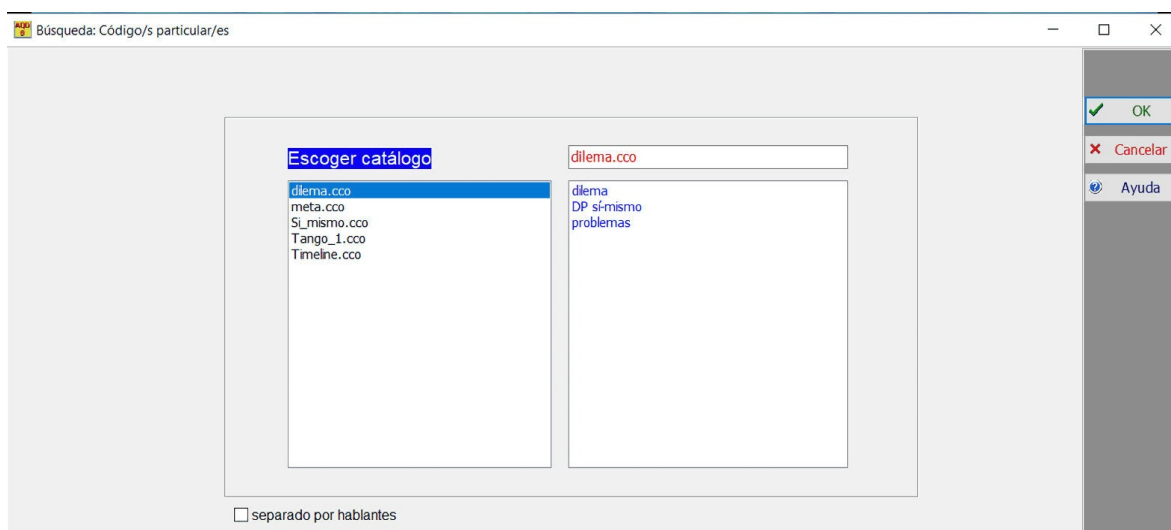
Hay muchas razones por las que, de vez en cuando, en el transcurso de un proyecto, un investigador puede querer hacer un seguimiento de todos los pasajes de texto que tratan el mismo tema, es decir, pasajes de texto que entran en la misma categoría. Una de estas razones, que desempeña un papel importante en muchos estudios descriptivos-interpretativos (véase Tesch, 1990), es detectar puntos comunes en toda la base de datos. Otra razón es el esfuerzo por controlar la coherencia de la codificación, es decir, garantizar que los códigos se han utilizado siempre según los mismos principios. Otras razones pueden encontrarse en la búsqueda de pruebas de dónde se ha utilizado correctamente un código y dónde se ha utilizado incorrectamente, dónde hay paralelismos entre ciertos textos, etc.

AQUAD extrae de todos los textos los segmentos que busca. Esto se hace en el orden en que los textos aparecen en el directorio de archivos. Dado que se busca un texto tras otro, también llamamos a esta búsqueda de segmentos de texto relevantes análisis unidimensional o lineal. Se diferencia de los análisis bidimensionales, llamados análisis de tablas o matrices en AQUAD, y de los análisis de enlaces complejos. Los resultados de los análisis unidimensionales pueden leerse en pantalla, imprimirse en papel o guardarse inicialmente en disco.

Búsqueda de códigos específicos

La búsqueda lineal de segmentos de texto codificados se inicia en el menú principal "Búsqueda" con la opción "código específico". Se abre una ventana en la que puede decidir qué buscar seleccionando uno de los catálogos de códigos disponibles.

En el ejemplo de la página siguiente, se ha trabajado en el proyecto "Entrevistas" que incluye los archivos "entrevista_1.txt" a "entrevista_4.txt". Para ello, entre otras cosas, se creó un catálogo con cuatro nombres de código, que contiene los códigos de perfil "/escuela primaria" y "/escuela secundaria", así como dos códigos conceptuales "dilema", "DP sí-mismo" y "problemas". Este catálogo "dilema.cco" ("cco" por "catálogo de códigos") ha sido seleccionado para el ejemplo haciendo clic en él:



Un clic en el botón "OK" devuelve inmediatamente el resultado de la búsqueda, pero a causa de espacio podemos mostrar aquí solamente los resultados de las dos primeras entrevistas del proyecto:

```
work}}} - Editor
Datei Bearbeiten Format Ansicht Hilfe
| Búsqueda de codificaciones específicas en >>entre_1.txt<<
-----
/--> dilema
    57 -    65: dilema
    136 -   141: dilema

/--> DP sí-mismo

/--> problemas
    51 -    65: problemas

3 hallazgo/s

Búsqueda de codificaciones específicas en >>entre_2.txt<<
-----
/--> dilema
     8 -    27: dilema
    110 -   129: dilema

/--> DP sí-mismo
    26 -    27: DP sí-mismo
    91 -    94: DP sí-mismo
    117 -   129: DP sí-mismo

/--> problemas
    34 -    34: problemas
    37 -    53: problemas
    147 -   150: problemas

8 hallazgo/s
```

No queremos discutir los hallazgos en este momento, sino que sólo nos preocupa la forma en que se presentan: Podemos ver en los archivos de código (*entre_1.aco*) que la entrevista 1 se realizó a un profesor de la escuela secundaria ("*escuela secundaria*") que reflexiona sobre los problemas en cuatro lugares, pero nunca sobre sí mismo. La entrevista 2 es con una profesora de la escuela primaria ("*escuela primaria*", en *entre_2.aco*) que habla de problemas en seis segmentos del texto y de reflexiones en cuatro lugares. El resultado aparece en el editor de texto y, por tanto, puede imprimirse, guardarse, copiarse en su totalidad o en extractos en el módulo "*Archivo*" (arriba a la derecha en la barra de menú del editor).

7.1.2 Cómo buscar estructuras de código

Búsqueda de estructuras de código: Códigos anidados

Esta opción busca si los subcódigos siguen encerrados dentro de los límites de las líneas de un código. O en otras palabras: la opción abarca todos los códigos que encierran otros códigos dentro del segmento de texto asignado. Por lo tanto, esta opción de menú también es adecuada para buscar secuencias de codificación estructuradas jerárquicamente. De este modo, también se muestran los códigos con líneas iniciales y/o finales idénticas, siempre que sus segmentos de texto no se superpongan (véase más adelante).

Ejemplo: En nuestro proyecto de ejemplo "Entrevista", algunos segmentos se codificaron con "problemas". Más tarde, al codificar, prestamos atención a los detalles y anotamos cuál era el problema, por ejemplo, dudas respecto a mantener el orden en el aula o favorecer la espontaneidad de los alumnos. Ahora queremos vincular los dos. Por ejemplo, una búsqueda podría mostrar que un segmento concreto del texto estaba codificado como "problemas", mientras que dentro de ese segmento se había asignado el código "orden-espontan." a una sección más pequeña. AQUAD proporciona información al respecto junto con el texto del segmento como sigue:

Búsqueda de codificaciones anidadas: Códigos inferiores en >>entre_1.txt<<

--> dilema

57- 65: dilema

57- 65: adaptación

136- 141: dilema

136- 141: orden-espontan.

--> DP sí-mismo

--> problemas

51- 65: problemas

51- 65: /\$Profesor

57- 65: adaptación

57- 65: dilema

--> orden-espontan.

136- 141: orden-espontan.

136- 141: dilema

--> indiv.-grupo

--> confianza-reglas

--> directivas-libertad

4 hallazgo/s

Vemos que en el primer texto se diagnostica problemas en varios lugares, una vez vinculado con "problemas" y además con "adaptación" (al nivel de los alumnos; líneas 57- 65).

Además, AQUAD puede revelar estructuras de codificación jerárquicas según dos direcciones:

- de los códigos de nivel superior, se buscan los códigos de nivel inferior (como en el ejemplo anterior);
- a partir de los códigos subordinados, se buscan los códigos de nivel superior, cuyos segmentos de texto encierran así los segmentos de los códigos subordinados.

Búsqueda de estructuras de código: códigos superpuestos

Como resultado de esta estrategia de búsqueda, se muestran todos los códigos (como siempre con información sobre la "ubicación") con los que se interpretaron los segmentos de texto superpuestos. Aquí se nos informa de los segmentos de texto en su conjunto, no sólo de la parte más estrecha en la que se solapan los códigos.

Ejemplo: Volvemos a tomar nuestros textos de ejemplo "entrevista" y buscamos los segmentos de texto que forman una intersección con los que fueron codificados con "IC colegas" (catálogo "Colegas.cco").

En los textos "entre_2.txt" y "entre_4.txt" este código no se utilizó nunca, mientras en el texto "entre_1.txt" vemos que se menciona los colegas una vez y en texto "entre_3.txt" dos veces (en relación con varios otros códigos).

Tenga en cuenta que el "solapamiento" no es más que una superposición física de pasajes de texto, sin que haya necesariamente referencias semánticas.

Otra nota sobre los códigos de perfil (véase más arriba): los códigos de perfil caracterizan todo el texto, por ejemplo el género del hablante ("/femenino") o el tipo de escuela donde el hablante trabaja ("/escuela primaria"). No obstante, deben insertarse en una sola línea, preferiblemente al principio del texto, ya que de lo contrario se solaparían con todos los códigos al buscar solapamientos sin aportar información relevante.

Búsqueda de estructuras de códigos: códigos múltiples

Esta estrategia de búsqueda se utiliza para encontrar todas las ubicaciones de los archivos que están asociadas a más de un código. Digamos que no es del todo seguro, siempre establecemos el código "problemas" al codificar las entrevistas a los profesores si el entrevistado habla de dificultades en un segmento de texto. Sin embargo, como no estábamos seguros de que ciertos focos de problemas no surgieran en entrevistas posteriores, intentamos marcar cada vez el mismo segmento de texto con un código que indicara el tipo de dificultad, por ejemplo, "problemas" y también, por ejemplo, "falta de conocimiento de los alumnos". Más adelante, nos damos cuenta de que, en lugar de hablar de "problemas" en general, deberíamos diferenciar mejor las áreas problemáticas específicas. Ahora buscamos qué dificultades notificadas hemos marcado con "problema" y con qué otros códigos al mismo tiempo:

Búsqueda de codificaciones múltiples en >>entre_1.txt<<

--> problemas

51 - 65: problemas |

51- 65: /\$Profesor

1 hallazgo/s

Búsqueda de codificaciones múltiples en >>entre_2.txt<<

--> problemas

```
34- 34: problemas |
    34- 34: /$Profesor
37- 53: problemas |
    37- 53: /$Profesor
```

2 hallazgo/s

Búsqueda de codificaciones múltiples en >>entre_3.txt<<

--> problemas

```
45- 54: problemas |
    45- 54: CA conocim. previos
    45- 54: clase - alumnos
    45- 54: Meta-CA
```

1 hallazgo/s

Búsqueda de codificaciones múltiples en >>entre_4.txt<<

--> problemas

```
80- 87: problemas |
    80- 87: /$Profesor
    80- 87: CE disciplina
```

1 hallazgo/s

Podemos ver que en las primeras dos entrevistas no hemos diferenciado los "problemas"; entonces deberíamos volver a los segmentos señalados y buscar más información.

Búsqueda de estructuras de código: secuencias de codificación

Un primer paso útil para abordar la reconstrucción de los contextos de significado en los textos es averiguar qué enunciados aparecen con frecuencia cerca de otros enunciados. Determinamos un código crítico como punto de referencia para la búsqueda de secuencias de códigos. AQUAD informa de todas las combinaciones de códigos de los segmentos del fichero que se producen a una distancia máxima (que definimos además) del código que es el punto de referencia. También se muestran los estados de "anidado" y "superpuesto".

Ejemplo: Queremos averiguar rápidamente qué otros códigos se encuentran en las proximidades de cada segmento de archivo codificado con "DP sí-mismo" (reflexión del profesor sobre sí mismo). Definimos como distancia máxima 3 líneas (unidades de distancia) antes o después de los segmentos críticos del archivo. Esto significa que AQUAD informa de todos los códigos asociados a los segmentos del archivo que terminan un máximo de 3 líneas antes del contenido marcado con el código crítico "DP sí-mismo" y de todos los códigos asociados a los segmentos del archivo que comienzan un máximo de 3 unidades de distancia después de la última línea/unidad de tiempo/imagen del segmento crítico (en la figura siguiente se puede ver sólo los resultados de "entre_2.txt"):

- Los códigos marcados con "<-" comienzan a una distancia de como máximo tres unidades (aquí: líneas) antes del segmento crítico (pero: también se señala la incrustación),
- Los códigos marcados con "~" están dentro del segmento crítico,
- Los códigos marcados con "->" indican segmentos que comienzan como máximo tres unidades después del segmento crítico.

Búsqueda de secuencias de códigos en >>entre_2.txt<<

Distancia máxima 3 unidades de distancia

--> DP sí-mismo

- 26 - 27: DP sí-mismo
 - <- 5 - 29: /\$Profesor
 - <- 8 - 27: confianza-reglas
 - <- 8 - 27: dilema
 - <- 8 - 27: orden-espontan.
 - <- 21 - 25: CE disciplina
- 91 - 94: DP sí-mismo
 - <- 87 - 89: DP opiniones
 - <- 87 - 106: /\$Profesor
 - ~~ 94 - 106: CE metodología
- 117 - 129: DP sí-mismo
 - <- 108 - 129: /\$Profesor
 - <- 110 - 129: dilema
 - <- 110 - 129: orden-espontan.
 - > 130 - 133: \$no contar
 - > 130 - 133: /\$Entrevistador

Una función importante de AQUAD es encontrar relaciones entre eventos o categorías significativas en los datos o comprobar si existen dichas conexiones. El programa utiliza una serie de algoritmos en el módulo "*Vínculos*" que se basan en el principio de la deducción hacia atrás. Esto significa que AQUAD comprueba las suposiciones sobre cómo podrían relacionarse los códigos entre sí y busca en todos los archivos para ver si los códigos críticos pueden encontrarse en determinadas secuencias y distancias.

Cada segmento de archivo codificado puede convertirse en evidencia de una categoría correspondiente dentro de su archivo. A medida que leas tus textos, escuches audios, veas vídeos, los estructures, los compares y eches un vistazo a tus memos de vez en cuando, desarrollarás suposiciones sobre las relaciones entre algunas categorías que utilizas en tu análisis. Estas suposiciones sobre las relaciones pueden conducir a hipótesis sobre el contenido latente de los archivos. La búsqueda de secuencias de codificación ayuda a generar dichas hipótesis.

7.1.4 Excursus: Distancia entre segmentos de ficheros

Para las hipótesis sobre los vínculos entre las codificaciones o los segmentos de archivos codificados, la distancia entre estos segmentos desempeña un papel importante. Quedémonos con los ejemplos de textos de entrevistas con los que hemos trabajado hasta ahora en este capítulo y supongamos que tenemos razones para creer que algunos hablantes empiezan a pensar en sí mismos cada vez que hablan de un problema en su práctica diaria. En concreto, esto significa que nos damos cuenta de que se habla de un "*problema*" en un determinado segmento de texto y ahora esperamos que (en el proyecto de ejemplo "*Entrevista*") le siga un segmento que se ha codificado con "*DP sí-mismo*".

Dado que el programa no puede realizar un análisis semántico de la conexión de contenidos, la única solución aproximada es limitarla especificando la distancia máxima permitida entre los segmentos del flujo de texto. Hablar de "uno mismo" cinco minutos (o tres páginas más adelante en la transcripción) después de que un entrevistado haya llegado a hablar de posibilidades de comportamiento contradictorias (por ejemplo, la espontaneidad de los alumnos frente a las normas en el aula) no estará probablemente relacionado directamente con el problema mencionado. Si es así, los interlocutores suelen recoger explícitamente un pensamiento pasado, y tendríamos que marcar esta referencia aquí de nuevo con el código "*problema*". Como valor por defecto, se especifican tres "unidades" en los análisis, aquí las líneas como la distancia máxima. Los segmentos de archivo que, con esta configuración, sólo

comienzan cuatro líneas después del final del primer segmento en la supuesta secuencia "Problemas" – "DP sí-mismo" no se consideran como hallazgos, aunque deban codificarse con "DP sí-mismo". Por lo tanto, debemos experimentar con la distancia máxima establecida para que se ajuste a las características de nuestros datos o comprobar el contexto en el texto original.

En el caso de los archivos de texto, el significado de la especificación de la distancia es claro: distancia máxima tolerada en número de líneas entre el final del primer segmento y el comienzo del segundo buscado. Pero, ¿qué significan las cifras para los archivos de audio o vídeo?

Archivos de audio:

El reproductor multimedia cuenta el flujo del audio en décimas de segundo. Para la comprobación de la distancia, el valor del contador se multiplica por un factor de 10, es decir, la distancia máxima se comprueba en unidades de segundos. La configuración por defecto $d=3$ tiene, por tanto, el efecto de que todos los segmentos de archivo codificados relevantes que comiencen un máximo de 3 segundos después del primer segmento encontrado se siguen aceptando como "vinculados" a él.

Archivos de vídeo:

El reproductor multimedia cuenta el flujo de vídeo en fotogramas. Para la comprobación de la distancia, el valor del contador se multiplica por un factor de 25, es decir, la distancia máxima se comprueba en unidades de segundos. La configuración por defecto $d=3$ tiene, por tanto, el efecto de que todos los segmentos de archivos codificados relevantes que comiencen un máximo de 3 segundos (o 75 fotogramas) después del primer segmento encontrado siguen siendo aceptados como "vinculados" a él.

7.1.5 Códigos NO utilizados

Esta opción del menú "Búsqueda" inicia una búsqueda unidimensional. A veces, en un proyecto grande, es útil saber qué códigos no se utilizan junto con qué archivos. AQUAD utiliza el contenido actual de su registro de códigos o una selección de códigos como criterio para esta búsqueda.

7.1.6 Contar códigos e introducir frecuencias en las tablas

Con la información sobre la frecuencia de los códigos seleccionados en un proyecto, son posibles muchos análisis cuantitativos que resultan bastante útiles, dependiendo de la pregunta de investigación. Se pueden crear tablas de frecuencias para exportarlas a programas de análisis cuantitativo (por ejemplo, hojas de cálculo y cálculos estadísticos) para todos los códigos, por ejemplo, también para los códigos de secuencia que normalmente sólo surgen en el curso de un análisis. En cualquier caso, el punto de partida es contar la frecuencia de los códigos con la función "Contar códigos" del submenú "Búsqueda". El resultado, una simple tabla de frecuencias de los códigos contados, se guarda automáticamente con el nombre del catálogo de códigos utilizado en el subdirectorio "..\res". En el caso del catálogo de códigos "Problema", encontramos la lista de frecuencias en este directorio como "Problema.txt".

Para el posterior tratamiento estadístico de la información sobre frecuencias, esta lista se convierte automáticamente en una tabla estructurada y se almacena también en el directorio "..\res". Las columnas están definidas por los códigos, las filas contienen los valores de frecuencia de estos códigos en los casos individuales (archivos). De

nuevo en el caso del catálogo de códigos "Problema" encontramos la tabla de frecuencias en este directorio como "Problema_DT.csv". La parte del nombre "_DT" significa "tabla de datos", la extensión ".csv" significa "valores separados por comas" (de hecho, el punto y coma se utiliza como separador). Estas tablas CSV tienen una estructura estandarizada para poder exportarlas sin problemas a programas de análisis de hojas de cálculo (por ejemplo, Microsoft Excel) o de análisis estadístico (por ejemplo, R o SPSS).

7.2 Cómo encontrar relaciones entre codificaciones

Ayudar a los investigadores a encontrar y comprobar las relaciones de significado en los datos es la función central de AQUAD. Para ello, el programa utiliza algoritmos que hacen realidad el principio lógico de la "deducción hacia atrás". Con ello, AQUAD busca en los archivos de código y comprueba si los elementos de una supuesta conexión pueden confirmarse en el orden y la distancia correctos. Según Glaser y Strauss, cada uno de los muchos códigos de un proyecto representa una categoría analítica: "Una categoría se sostiene por sí misma como elemento conceptual de la teoría" (Glaser & Strauss 1967, p. 36). Cada codificación o cada segmento de datos marcado con un código es una prueba de la aparición de la categoría correspondiente en los datos.

A medida que lea, organice y compare los datos y consulte la lista de códigos y quizás los memorandos relacionados, es probable que se le ocurran conjeturas al principio del proceso de análisis sobre las relaciones que parecen existir entre algunas de las categorías de su análisis. "Cada vez que estas personas hablan de acontecimientos vitales críticos, te surgen emociones...", puedes notar. Estos vínculos pueden convertirse en la ocasión para formular hipótesis o teorías iniciales sobre lo que los datos realmente tratan.

Ya se ha señalado anteriormente que en la práctica siempre se combinan los movimientos de pensamiento inductivo y deductivo. Por ejemplo, también se puede partir de hipótesis sobre posibles correlaciones a la inversa de lo que se acaba de describir e intentar demostrarlas a partir de los datos, es decir, deducir pruebas de estas hipótesis a partir de los datos. Si no se encuentra nada que pueda confirmar las suposiciones o hipótesis generales sobre la conexión entre las categorías, o si se encuentran contradicciones, entonces se realiza un esfuerzo inductivo para desarrollar categorías y sus conexiones sistemáticas a partir de los datos, es decir, para generalizar conexiones básicas de significado a partir de segmentos de datos llamativos. Esta es, por supuesto, una representación extremadamente abreviada del procedimiento inductivo-deductivo en el análisis de datos cualitativos.

En cualquier caso, sin embargo, se llega al punto de formular una suposición o vinculación de categorías generada lentamente y buscada desde el principio como una afirmación sobre precisamente esta conexión. En otras palabras, se formula una hipótesis. En el análisis de datos cualitativos, una hipótesis de este tipo es básicamente una afirmación de que determinados segmentos de datos codificados o las categorías que deben reconocerse en ellos se dan de forma combinada. Un ejemplo podría ser la conjetura que surgió en el análisis de las entrevistas biográficas: "Existe un vínculo entre los acontecimientos vitales críticos y las reacciones emocionales".

En principio, comprobar cómo se relacionan los códigos es muy sencillo. Usted introduce los códigos que cree que están relacionados de una manera determinada y entonces el ordenador empieza a buscar si hay segmentos en sus archivos que estén codificados de esa manera. Es decir, el ordenador busca códigos que suponen que están específicamente vinculados y el ordenador busca características que son indicativas de tales vínculos. La ventaja del ordenador es que no se pierde nada. Se puede confiar en ella para encontrar y considerar cualquier codificación relevante en los datos.

Para adaptar este principio de búsqueda a sus necesidades, AQUAD ofrece diversas variantes. De forma más general, hay cuatro formas de comprobar los vínculos entre los códigos o de poner a prueba sus hipótesis. Estas cuatro posibilidades se presentan a continuación según las diferentes estrategias de investigación. AQUAD busca

- conexiones en el contexto de las categorías incondicionales, es decir, en un intervalo de líneas que debe definirse antes y/o después de los segmentos de texto de una categoría crítica. Dentro de esta área de texto, se registra la aparición de cualquier otra codificación sin tener que cumplir ninguna otra condición adicional.
- contextos en el marco de las categorías condicionales. La "condición" para la posible existencia de correlaciones sistemáticas está definida por una segunda categoría a la que debe asignarse un segmento de datos al mismo tiempo.
- conexiones en forma de secuencias de codificación simples. Estas secuencias se han incorporado a AQUAD durante su desarrollo con motivo de investigaciones concretas. Basta con insertar los códigos en las estructuras de enlace abstractas.
- interrelaciones en forma de relaciones complejas de codificación. Para la verificación, puede poner las relaciones en forma de hipótesis comprobables sobre las secuencias de codificación que sospecha en sus datos.

7.2.1 Relaciones en el contexto de las categorías incondicionales: Búsqueda de códigos

Ya se ha familiarizado con esta forma más sencilla de búsqueda o comprobación no lineal en los textos de datos en la presentación de las distintas posibilidades de recuperación de segmentos de datos codificados o de los códigos utilizados en ellos con la ayuda de estructuras de búsqueda generales o específicas.

El ordenador sólo se encarga de registrar todos los códigos a una determinada distancia alrededor del código que acaba de encontrar ("análisis de distancia"). Algo más compleja es la orden de buscar patrones de codificación específicos, por ejemplo, segmentos de datos codificados múltiples. En el caso de la búsqueda específica, se sigue imponiendo la restricción de que sólo se comuniquen los patrones en los que esté representado un código específico. La "categoría" en el caso general es, por tanto, una estructura de codificación específica.

Lo que no se requiere aquí para el éxito de la búsqueda es que el segmento de datos pertenezca al ámbito de (al menos) otra categoría. Por ejemplo, los algoritmos de búsqueda descritos en el capítulo 7.1.2 no pudieron distinguir entre los solapamientos de otros segmentos de datos con el ámbito de validez de, por ejemplo, la categoría "acontecimiento vital crítico" en las entrevistas de "personas jóvenes" y "personas mayores". Se trataría de una búsqueda de categorías "condicionadas" (aquí condicionadas por la característica "edad").

7.2.2 Correlaciones en el contexto de las categorías condicionales: Análisis de tablas o matrices

Miles y Huberman (1994) recomiendan y utilizan en su libro la forma de representación matricial como la estrategia más importante para estructurar el contenido textual. En particular, es importante para ellos hacer más manejable la configuración secuencial de enunciados, pensamientos y opiniones en los textos, transformándolos en configuraciones simultáneas. Según Miles y Huberman (1994, p. 93 y siguientes), las matrices o tablas ayudan a

- obtener o mantener una visión de conjunto, ya que los datos y los análisis se presentan de forma resumida;
- para reconocer rápidamente los casos en los que se necesitan análisis adicionales y más diferenciados;

- comparar datos e interpretaciones;
- comunicar los resultados de la investigación a otros y hacerlos comprensibles.

En Miles y Huberman (1994, capítulos IV, V y VI) se pueden encontrar más orientaciones sobre la organización de los datos verbales en tablas o matrices, junto con numerosos ejemplos. Para la recuperación asistida por ordenador, construya matrices de codificación (Miles y Huberman, 1994) o tablas de codificación (Shelly y Sibert, 1992). La doble determinación del contenido (segmentos de datos) de las celdas de una tabla de este tipo determina la interpretación de los resultados con más fuerza que la mera constatación de una acumulación de proximidad espacio-temporal de categorías inicialmente independientes. Por otro lado, para la construcción de una matriz de análisis es necesario un mayor trabajo conceptual previo y, por tanto, un mayor avance en el proceso de análisis.

Ejemplo: En el análisis de las entrevistas con los profesores, comparamos las entrevistas con los profesores de primaria y de secundaria para demostrar el análisis de la tabla. Con el nivel escolar como característica "singular", marcamos las columnas de nuestra tabla. Para ello, introducimos los códigos de perfil *"escuela primaria"* y *"escuela secundaria"*. Queremos poner a prueba la hipótesis de que existen diferencias entre los profesores y las profesoras con respecto a la infraestructura del centro escolar (*"IC infraestructura"*), los problemas en el aula (*"problemas"*), su propio aprendizaje profesional (*"DP aprender/prof."*) y las reflexiones sobre la propia persona (*"DP sí-mismo"*). (Por supuesto, para ello tendríamos que realizar más entrevistas seleccionadas sistemáticamente; pero para la presentación del principio, éstas serán suficientes aquí). Los códigos mencionados arriba definen las filas de la tabla. Al iniciar el análisis de la tabla, AQUAD recogerá todos los segmentos del archivo en los que los profesores hablan de aprendizaje profesional; en la columna adyacente se leen las declaraciones de las profesoras sobre el mismo tema. En la fila de abajo de la tabla encontramos primero las declaraciones de los profesores de primaria sobre sus reflexiones personales, en la columna de al lado las correspondientes declaraciones de las profesoras, y así sucesivamente.

Estas matrices son especialmente propicias para las exigencias de la estructuración de los datos, ya que estructuran y presentan una amplia información de forma simultánea (en lugar de secuencialmente como en los archivos originales), permiten comparaciones rápidas con los resultados de otros sujetos, situaciones, puntos temporales, investigaciones, etc., y abren nuevos análisis más refinados. Si es cierto que una imagen vale más que mil palabras, una presentación tabular sustituye al menos a cien palabras, sobre todo porque permite visualizar los datos y el proceso de análisis de forma conjunta.

Sin embargo, aquí parece apropiada una advertencia: incluso una presentación tabular muy diferenciada de los datos reducidos no puede hacer más que proporcionar buenas condiciones previas para otras conclusiones, pero no puede en absoluto sustituir las conclusiones necesarias. La representación matricial también alcanza sus límites cuando no sólo se utilizan características singulares como la infraestructura escolar o el sexo como criterios de ordenación, sino cuando se buscan combinaciones específicas de unidades de significado, cada una de las cuales puede aparecer varias veces en un fichero. Para estas evaluaciones de orden superior, hay que especificar patrones condicionales complejos, que suelen comprender varios códigos o unidades de significado y sus relaciones lógicas, y que a veces también incluyen datos cuantitativos, por ejemplo, sociodemográficos. Una matriz bidimensional se vería desbordada por las exigencias de estos análisis; una disposición de matrices más diferenciadas o especializadas se volvería rápidamente inmanejable.

Las distintas funciones del análisis de tablas se encuentran en la opción del menú principal *"Tablas"*. Debe estudiar detenidamente el uso de las funciones *"Crear tabla"* y *"Editar tabla"*. Las tres opciones siguientes, *"Análisis: frecuencias"*, *"Análisis: codificación"* y *"Análisis: segmentos de texto"*, determinan la forma en que se obtienen los resultados. Existen tres opciones: indicar simplemente la frecuencia de las codificaciones, indicar adicionalmente las

ubicaciones de las codificaciones o incluso una impresión completa de los segmentos de texto codificados que cumplen las condiciones de la construcción de la tabla.

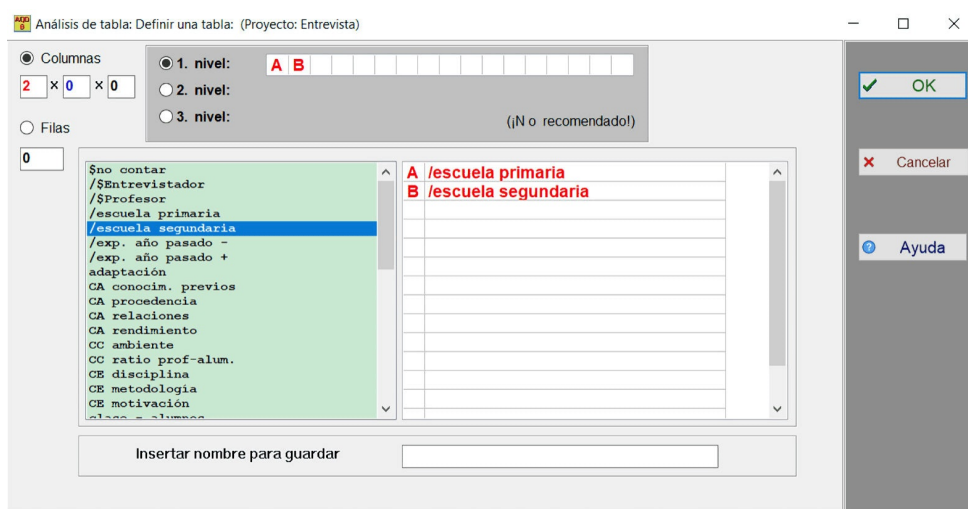
Cómo construir una tabla

Para construir una tabla para una búsqueda significativa de datos bidimensionales en AQUAD, se necesitan al menos dos archivos que difieran en una característica del perfil (o característica "singular"). Sólo se pueden buscar segmentos o códigos con una matriz o determinar su frecuencia si los segmentos del archivo se han interpretado con al menos dos códigos diferentes.

Uno de estos códigos debe ser un código sociodemográfico o de perfil de todo el texto, es decir, un código singular que se asigna una sola vez en un texto y que caracteriza a todo el texto. En el ejemplo del análisis de la entrevista, se ha introducido una categoría cuyas características podrían representarse cada una con un código de perfil: *"/escuela primaria"* y *"/escuela secundaria"*. Ambas son características del interlocutor de la entrevista o de su texto en su conjunto y pueden utilizarse para definir las columnas:

Los demás códigos que definen las filas de la tabla deben pertenecer a los códigos conceptuales del estudio. Pueden y probablemente ocurrirán más de una vez por archivo de código. En el ejemplo de la entrevista, seleccionamos los códigos infraestructura del centro escolar (*"IC infraestructura"*), los problemas en el aula (*"problemas"*), el propio aprendizaje profesional del profesor/de la profesora (*"DP aprender/prof."*) y las reflexiones sobre la propia persona (*"DP sí-mismo"*) como interesantes para el análisis:

Ahora, cada vez que AQUAD encuentre uno de estos códigos en un archivo de códigos, (si la opción "Análisis: Segmentos de texto" está seleccionada) también emitirá el segmento de texto asociado. Sin embargo, primero debemos averiguar cómo se introducen estos códigos: Tras activar la opción *"Crear tabla"*, aparece la siguiente ventana. Verá dos botones redondos en la esquina superior izquierda: *"Columnas"* y *"Filas"*.



La captura de pantalla muestra que la opción *"Columnas"* ya está activada (hay un pequeño punto en el botón) y que los códigos de perfil *"/Escuela Primaria"* y *"/Escuela secundaria"* se han establecido como encabezados de columna (haciendo clic en ellos en el registro (verde) de códigos).

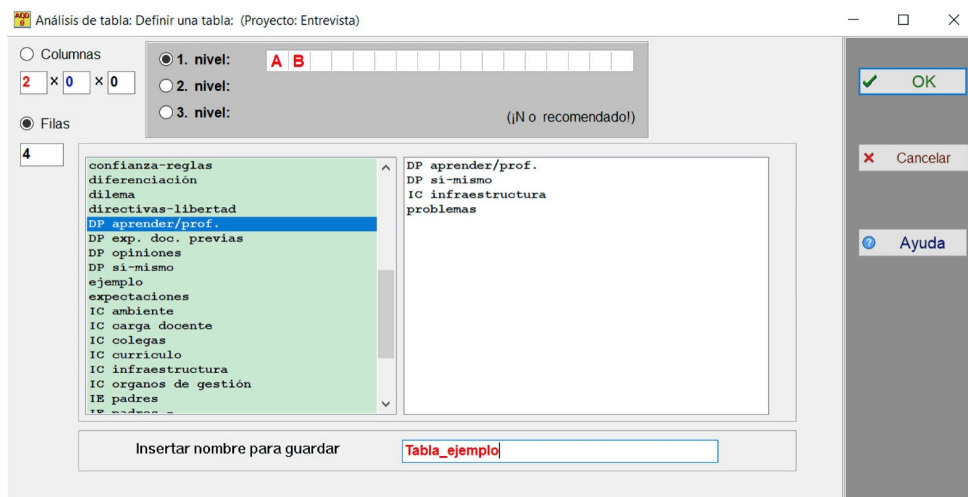
Por favor, tenga en cuenta:

Sólo puede utilizar códigos de perfil como encabezamiento de columna (indicador: "/"; aquí también se denominan códigos singulares).

En nuestro proyecto de ejemplo "Entrevista" hemos hecho clic en los códigos de perfil *"Escuela primaria"* y *"Escuela secundaria"*. A continuación, ambos se transfirieron al cuadro blanco de la derecha y las abreviaturas "A" y "B" se introducen automáticamente en las pequeñas celdas que siguen detrás de la etiqueta "1.nivel". Como puede deducirse del número de celdas, es posible un máximo de 12 encabezados de columna en el primer nivel. Al mismo tiempo, las tres primeras casillas de recuento de la izquierda se rellenan automáticamente con números: primero aparece "1" y luego "2": se cuentan los títulos de las columnas (= condiciones de análisis).

Si tenemos muchos archivos de texto, podemos diferenciar los títulos de las columnas en un máximo de seis títulos de "segundo nivel". Así, en el ejemplo del análisis de las entrevistas, podríamos subdividir adicionalmente las categorías de nivel escolar según el género o la experiencia profesional (principiante, experiencia práctica, experto) de los profesores. De esta manera creamos columnas de 2x2 o 2x3. Sin embargo, para un análisis de este tipo necesitaríamos un mayor número de archivos que puedan asignarse a los determinantes de las columnas. En el proyecto de ejemplo *"Entrevista"* sólo tenemos cuatro textos o casos, por lo que la diferenciación no tendría mucho sentido aquí. A veces, cuando se dispone de muchos textos, se puede pensar incluso en un tercer nivel de diferenciación – en principio, esto es posible en AQUAD, pero no suele ser muy aconsejable por la dificultad de interpretar los resultados.

Después de establecer los títulos de las columnas, es el turno de las filas. Haga clic en el botón "Filas" y, a continuación, seleccione los códigos conceptuales adecuados, por ejemplo, *"IC infraestructura"*, *"problemas"*, etc., como se ha sugerido anteriormente.



La selección se transfiere de nuevo automáticamente al cuadro blanco de la derecha y el número de líneas previstas se cuenta en el pequeño cuadro situado debajo del botón "Filas".

Ahora estamos casi listos para pulsar el botón "OK". Sin embargo, en la parte inferior de la ventana hay un pequeño cuadro en el que se nos pide que introduzcamos un nombre (¡sin extensión!) para guardar la definición de nuestra tabla. Tras hacer clic en "OK", todas las entradas se guardarán para un posterior análisis de la tabla.

Le sugerimos que pruebe todo esto una vez según la descripción arriba. Si quieres ver el aspecto de la ventana después de haber etiquetado columnas y filas, ve primero a "*Editar tabla*" (en el submenú "*Tablas*") y selecciona la definición de la tabla que has creado y – ojalla – guardado. Salte de las opciones "*Columnas*" y "*Filas*" para echar un vistazo.

Resumen:

1. Haga clic en el botón "*Columnas*" y seleccione el primer nivel de encabezamientos de columna.
2. Se introducen hasta doce códigos sociodemográficos (de perfil) como títulos de columna.
3. Si es necesario para el análisis, añada otro nivel con los correspondientes títulos de columna (códigos de perfil subordinados).
4. Haga clic en "*Filas*" y seleccione los códigos conceptuales como títulos de las filas.
5. Introduzca un nombre para guardar la definición de esta tabla y haga clic en "*OK*".

Cómo editar las tablas

En principio, se aplica la misma secuencia de pasos descrita anteriormente para la elaboración de una tabla o matriz. Sin embargo, primero debe decidir cuál de las definiciones de tabla ya creadas quiere editar.

Si no quiere sobrescribir el diseño de la tabla que acaba de editar, introduzca un nuevo nombre para guardarlo en el pequeño campo de entrada que aparece a continuación (antes de pulsar "*OK*").

Cómo realizar un análisis de la tabla

Tiene tres opciones, que en principio funcionan igual, pero dan resultados informativos diferentes.

- El más económico, pero también el menos informativo, es el "*Análisis: Frecuencias*". Para ello, el resultado cabe en la pantalla o en una hoja de papel, y vemos de un vistazo la frecuencia con la que los códigos conceptuales aparecen en los textos bajo la condición de los códigos del perfil. Sin embargo, el lugar exacto y los enunciados de esta función permanecen ocultos para nosotros. Por ello, el análisis de frecuencias se utiliza principalmente como heurístico en la búsqueda de contextos. Veamos qué tipo de tabla de frecuencias proporciona la definición de tabla creada anteriormente en nuestro proyecto de ejemplo "Entrevista" (Si es necesario, la tabla de resultados puede guardarse en formato CSV como se ha descrito anteriormente):

The screenshot shows a window titled "Análisis de tabla: (Proyecto: Entrevista) frecuencias". Inside, there is a table with two columns, A and B, and five rows of data. Below the table, there are two input fields labeled A: and B: with the text "/escuela primaria" and "/escuela secundaria" respectively. To the right of the table, there are two buttons: "Cerrar" (Close) and "Ayuda" (Help). A small dialog box titled "Atención:" is open in the foreground, asking "¿Quiere guardar esta tabla?" (Do you want to save this table?) with "SI" (Yes) and "NO" (No) buttons.

	A	B
DP aprender/prof.	9	0
DP sí-mismo	9	3
IC infraestructura	1	5
problemas	5	7
A: /escuela primaria		
B: /escuela secundaria		

- Más detallado es el "*Análisis: Códigos*". Proporciona una lista en la que se introducen los códigos que allí se encuentran (es decir, los números de línea y los nombres de los códigos) para cada celda de la definición de la tabla, sucesivamente.
- Los resultados son más detallados con la opción "*Análisis: Segmentos de texto*". Esta opción devuelve todos los segmentos de texto que satisfacen las condiciones duales del diseño de la matriz. Por lo tanto, tenga en cuenta que si desea imprimir los resultados en la impresora, el consumo de papel puede ser considerable. Ni la impresora ni la pantalla pueden mostrar las columnas de la tabla llena de documentos de texto uno al lado del otro. Por lo tanto, AQUAD imprime una celda tras otra para cada columna. A continuación, puedes montar las impresiones en papel en una tabla sobre un tablero grande o cualquier otra superficie adecuada si quieres ver el resultado en su totalidad.

7.3 Establecer conexiones en forma de secuencias de codificación sencillas: Comprobación de los vínculos

Para los análisis que generan o construyen teorías, el módulo de programa "*Vínculos*" es probablemente el más importante del paquete de programas AQUAD. Utiliza la lógica deductiva para comprobar las conexiones hipotéticas entre las unidades de significado de los archivos. Cada código se considera un representante de una "categoría de análisis" en el sentido de Glaser y Strauss: "Una categoría se sostiene por sí misma como elemento conceptual de la teoría" (Glaser & Strauss, 1967, p. 36). Cada segmento codificado es la expresión actual de la aparición de esa categoría en sus datos. A medida que se trabaja con los datos y se observan las categorías conceptuales, surgen suposiciones sobre las relaciones que parecen existir entre las categorías. Estas relaciones y vínculos podrían ser el comienzo de observaciones teóricas sobre lo que ocurre en sus datos, o en las mentes de las personas de las que proceden los datos. Por lo tanto, también se puede decir que este módulo permite comprobar las hipótesis sobre las relaciones.

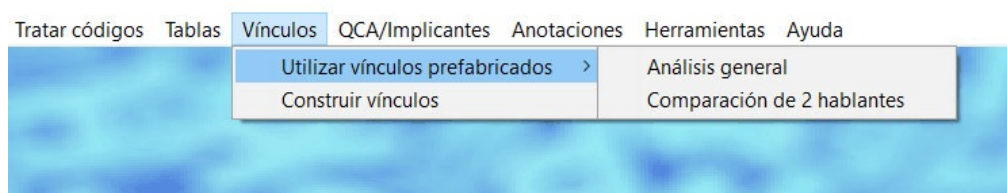
Se formula una presunta conexión entre las unidades de significado de un archivo como una afirmación de que esta conexión es verdadera y luego se deja que el ordenador compruebe todas las entradas de los archivos de codificación para ver si se cumplen las condiciones contenidas en la afirmación. Como resultado, no sólo recibirá un mensaje de que la prueba ha tenido éxito o ha fracasado, sino también una lista de todas las "pruebas", es decir, todas las codificaciones que corresponden a la afirmación. Esto puede servir para realizar otras comparaciones.

Dentro de AQUAD, se pueden comprobar todos los tipos de secuencias de codificación de esta manera. Hay algunos patrones de secuencia incluidos en el software para su uso inmediato como estructuras de enlace abstractas. Cuando se realiza un análisis de vinculación, se introducen algunos códigos concretos de su estudio, es decir, cuando se llaman, estas estructuras abstractas se convierten con los códigos dados de su propio estudio en formulaciones de hipótesis concretas que AQUAD puede cotejar con los archivos de codificación.

Para iniciar el análisis de vínculos, seleccione la opción "*Vínculos*" en el menú principal y luego "*Utilizar vínculos prefabricados*". Hay dos alternativas: se utiliza la opción "*Análisis general*" para buscar las secuencias de codificación generales que aparecen en todo el archivo dentro de los archivos de un proyecto. Sin embargo, esto le da la oportunidad de buscar los enlaces por separado para diferentes "hablantes" (generalmente: segmentos de archivos marcados por "códigos de hablantes") si es necesario. La opción "*Comparación de 2 hablantes*", en cambio, analiza las conexiones específicas en cada caso entre los segmentos de datos sucesivos de dos "hablantes" definidos (o, por ejemplo, también las preguntas de un cuestionario), por ejemplo: "Si el hablante 1 dice 'A', el hablante 2 afirma 'B'".

También se pueden comprobar correlaciones más complejas con AQUAD; sin embargo, debido a la variedad de correlaciones concebibles, no se incorporan al programa estructuras de secuencias abstractas para este fin. Para ello, la opción "*Construir vínculos*" le da la oportunidad de introducir sus suposiciones sobre enlaces complejos de códigos haciendo clic en los nombres de los códigos y en las conjunciones lógicas (operadores "AND", "OR" y "NOT"). A continuación se explica detalladamente cómo hacerlo. Por razones pragmáticas, el número de códigos que pueden vincularse de esta manera se limita a cinco.

El siguiente extracto del menú principal muestra de nuevo las posibilidades que ofrece AQUAD para buscar o confirmar vínculos:



Si su problema no puede resolverse también con las estructuras de vínculos o las posibilidades de construcción existentes, el autor estará encantado de crear un algoritmo a medida.

7.3.1 ¿Qué estructuras de vínculo proporciona AQUAD?

Secuencias de codificación generales / ejemplos

AQUAD proporciona 12 "vínculos" hipotéticos en los que, por lo general, puedes colocar cualquiera de tus códigos. Estas estructuras de enlace representan relaciones expresadas mediante fórmulas entre unidades de significado que se quieren comprobar.

Para ello, no sólo hay que especificar los códigos que se van a comprobar, sino también, por lo general, una distancia crítica que no debe superarse entre los segmentos de archivo codificados.

Estas estructuras no se han construido didácticamente, sino que se han desarrollado durante el trabajo con AQUAD en nuestras propias investigaciones, en la supervisión de tesis y en los cursos de doctorado que introducen métodos cualitativos. Antes de profundizar en estos vínculos y examinar algunos ejemplos, describimos lo que significan estas formulaciones muy abreviadas en la pantalla (véase la figura). La primera frase repite la hipótesis de enlace, la segunda frase describe lo que el ordenador comprueba cuando se activa esta hipótesis de vínculo en particular:

Seleccionar un tipo de vínculo:

- ☐ 1. (Código 1 Y código 2) en distancia determinada, solo casos positivos
- ☐ 2. (Código 1 Y código 2) en distancia determinada, casos positivos y negativos
- ☐ 3. (Cód. 1 Y Cód. 2 Y Cód. 3), distancias determinadas entre 1/2 y 2/3
- ☐ 4. ((Cód. 1 O Cód. 2) Y Cód. 3) -> Cód. 3 en distancia determinada
- ☐ 5. ((C1 O C2 O C3) Y C4) -> Cód. 4 en distancia determinada
- ☐ 6. ((C1 Y C2) Y C3) -> 1/2 en dist. def., C3 en cualquiera dist.
- ☐ 7. ((C1 Y C2) Y C3 Y C4) -> 1/2 en dist. def., 3 y 4 en cualquiera dist.
- ☐ 8. ((Cód. 1 Y (Cód. 2 Y Cód. 3) -> C2, C3 dentro segmento de C1)
- ☐ 9. ((C1 AND C2) AND C3) -> C2 dentro el segmento de C1, C3 en dist. def. de C1
- ☐ 10. ((Cód. hablante C1 Y C2 (Y C3 Y C4)) -> 3, 4 en distancia def. de 2))
- ☐ 11. ((C1 Y C2) O (C3 Y C4)) -> distancia def. entre 1 y 2 / 3 y 4
- ☐ 12. ((C1 O C2) Y (C3 O C4)) -> distancia def. entre 1/2 y 3/4

1. Dos códigos concurren en el mismo documento de datos dentro de una distancia especificada ¿Verdadero para qué casos?
2. Dos códigos concurren en el mismo documento de datos dentro de una distancia especificada ¿Verdadero o falso para qué casos?
3. Tres códigos concurren en el mismo documento de datos; el código #2 dentro de una distancia especificada respecto al código #1, y el código #3 dentro de una distancia especificada diferente respecto al código #1. ¿Verdadero para qué casos?
4. Uno o los dos de dos códigos concurren en el mismo documento de datos junto con un tercer código, que aparece dentro de una distancia especificada. ¿Verdadero para qué casos?
5. Uno, dos, o los tres de tres códigos concurren en el mismo documento de datos junto con un cuarto código, que aparece dentro de una distancia especificada. ¿Verdadero para qué casos?
6. Dos códigos concurren dentro de una distancia especificada en un documento que contiene un tercer código (en cualquiera distancia, o sea como condición general). ¿Verdadero para qué casos?
7. Dos códigos concurren dentro de una distancia especificada en un documento que contiene un tercer y cuarto código (en cualquiera distancia, o sea como condiciones generales). ¿Verdadero para qué casos?
8. Tres códigos concurren en el mismo documento de datos, con los códigos #2 y #3 como subcódigos del código #1. ¿Verdadero para qué casos?
9. Dos códigos ocurren en el mismo documento de datos, con el código #2 como subcódigo de código #1, y un tercer código aparece a una distancia especificada del código #1. ¿Verdadero para qué casos?
10. Dentro del segmento de datos de un hablante distinto aparecen un, dos o tres códigos en una distancia especificada. ¿Verdadero para qué casos?
11. Dos códigos determinados o dos otros códigos concurren dentro de una distancia especificada en un documento de datos. ¿Verdadero para qué casos?
12. Un código o un código alternativo concurre en un documento dentro de una distancia especificada de un segundo código o un código alternativo. ¿Verdadero para qué casos?

Si quiere trabajar con una de estas estructuras de enlace, seleccione la formulación correspondiente y, a continuación, introduzca los códigos y, en la mayoría de los casos, también las distancias máximas. Suponiendo que esté interesado en la hipótesis 3, en este caso se abre una ventana de entrada especial:

Utilizar vínculos prefabricados: Análisis general (Proyecto: Entrevista)

\$no contar
/Entrevistador
/Profesor
/escuela primaria
/escuela secundaria
/exp. año pasado -
/exp. año pasado +
adaptación
CA conocim. previos
CA procedencia
CA relaciones
CA rendimiento
CC ambiente
CC ratio prof-alum.
CE disciplina
CE metodología
CE motivación
clase - alumnos
confianza-reglas
diferenciación
dilema
directivas-libertad
DP aprender/prof.
DP exp. doc. previas
DP opiniones
DP si mismo
ejemplo
expectaciones
experiencia
IC ambiente
IC carga docente
IC colegas
IC currículo
IC infraestructura
IC organos de gestión
IE padres

Y

Distancia: 3

Continuar

Cancelar

Ayuda

Código de secuencia:

En esta captura de pantalla, tanto las distancias entre los segmentos de texto vinculados al Código 1 y al Código 2 como las distancias entre los segmentos vinculados al Código 2 y al Código 3 ya han sido introducidas. Verá en las pequeñas casillas de la derecha una distancia preestablecida de tres líneas como máximo. Sólo tienes que hacer clic en las casillas e introducir una distancia diferente si tiene sentido para tu proyecto. Ahora tenemos que introducir los nombres en clave. Para ello, haga clic en los códigos correspondientes en el registro de códigos de la izquierda. Se transfieren automáticamente al siguiente campo libre en el centro. Si es necesario, haga doble clic para vaciar estos campos de nuevo y transfiera el contenido de nuevo al registro de códigos si quiere formular la hipótesis de vínculo de forma diferente.

Si es necesario introducir una distancia en una fórmula de vínculo, puede eludir esta especificación introduciendo un número que corresponda al número total de líneas de su texto más largo - en duda: 99999. Con este truco, se registra cualquier ocurrencia de dos códigos.

También hay que tener en cuenta que el orden en que se introducen los códigos es crucial para el resultado, ya que la estructura de vínculo es "dirigida", es decir, determina la secuencia de los segmentos codificados. El primer código introducido debe aparecer también en primer lugar en el archivo. Si el segundo aparece primero y fuera de la distancia determinada, en este caso particular se rechaza como falsa la hipótesis de una vinculación del código 1 y el código 2 a una distancia definida.

En la parte inferior de la ventana encontrará la etiqueta "Código de secuencia" y detrás de ella un campo de entrada. ¿De qué se trata todo esto? Los resultados de los vínculos de códigos que son relevantes para la cuestión de su investigación pueden resumirse automáticamente en un único código, un "código de secuencia", que se extiende para el segmento de archivo desde el principio del primer código hasta el final del último código enlazado. Si queremos comprobar la siguiente hipótesis

En este proyecto hay entrevistas con profesores cuyas experiencias de problemas de enseñanza están vinculadas a reflexiones sobre sí mismos.

Por ejemplo, utilizaríamos la hipótesis de vinculación 2 ("*Dos códigos aparecen en el mismo archivo a una cierta distancia el uno del otro*") para buscar vínculos entre "*problemas*" y "*DP sí-mismo*" (distancia: 5 filas) en el proyecto "*entrevista*". Esta prueba de hipótesis arroja el siguiente resultado:

Estructura prefabricada de vínculo:
problemas Y DP sí-mismo
Distancia 5

```
=====
--> Archivo: entre_1.txt-----
  51- 65: problemas
0 Confirmación/iones
--> Archivo: entre_2.txt-----
  34- 34: problemas
  37- 53: problemas
 147- 150: problemas
0 Confirmación/iones
--> Archivo: entre_3.txt-----
   9- 18: problemas
  34- 37: problemas
   Y 34- 37: DP sí-mismo
  45- 54: problemas
  76- 82: problemas
 136- 137: problemas
 144- 149: problemas
1 Confirmación/iones
--> Archivo: entre_4.txt-----
  80- 87: problemas
 142- 153: problemas
   Y 145- 149: DP sí-mismo
   Y 150- 152: DP sí-mismo
2 Confirmación/iones
```

La lista de referencias muestra que
en la entrevista 1 y 2 el presunto vínculo nunca se produce,
en la entrevista 3 sólo una vez y
dos veces en la entrevista 4.

En caso de que la vinculación examinada arroje resultados significativos para nuestra investigación, podríamos introducir "*Reflexión en caso de problemas*" como código de secuencia en el campo de entrada de abajo. Tras hacer clic en "*Continuar*", aparece una pregunta de seguridad: "*¿Desea insertar los enlaces encontrados como códigos de secuencia en los archivos de códigos?*" Si responde "*Sí*" aquí, los archivos de códigos se ampliarán en consecuencia, de lo contrario sólo recibirá la lista de resultados.

Para cada búsqueda de vínculos, AQUAD ofrece la posibilidad de marcar los sitios de hallazgo con un código de secuencia, es decir, también de marcar los hallazgos con hipótesis de vínculo autoconstruidas. En la ventana correspondiente encontrará la máscara de entrada etiquetada como "*Código de secuencia*".

Comparación de la codificación entre dos oradores / ejemplos

Un tipo de códigos en AQUAD son los códigos de hablante ("/\$...."). Se aplican códigos de hablante para atribuir segmentos de datos – por ejemplo segmentos de grabaciones de discusiones en grupos pequeños – a los hablantes distintos. Otra aplicación sería codificar las contestaciones a preguntas abiertas en un cuestionario con "códigos de

hablante", en esta caso de verdad códigos de pregunta: /\$pregunta_1, /\$pregunta_2, etc.; por eso se podría analizar separadamente y comparar muy sencillamente las preguntas de los sujetos un sus archivos de dato diferentes.

La opción "*Utilizar vínculos prefabricados*" -> "*Comparar 2 hablantes*" asiste a buscar asociaciones específicas en segmentos de dato consecutivos de dos "hablantes" determinados (o también dos preguntas en un cuestionario). Preguntas de investigación típicas que se podría contestar con esta opción serían por ejemplo:

- ¿Es de verdad que alumnos contestan por frases fragmentarias cuando sus maestros hacen preguntas cerradas?
- ¿Es de verdad que una persona afirma cada vez el contrario cuando una otra persona ha dicho algo?
- ¿Es de verdad que sigue frecuentemente a la respuesta X a pregunta 3 de un cuestionario distinto la respuesta Y a la pregunta 7?

Preguntas como estas se podrían contestar sencillamente por las hipótesis de vínculos entre dos hablantes. A diferencia de vínculos generales (vea arriba) sirve como criterio de comprobación el hecho de que códigos determinados concurren en los segmentos de dato de hablantes (o preguntas de cuestionario) diferentes.

Las denominaciones A, B, C, ... de códigos en las descripciones siguientes no señalan más que buscamos segmentos de datos concretos determinados en las secciones de hablantes diferentes – no es necesario que tienen sentido diferente en las secciones diferentes. De otras palabras, según hipótesis 1 (vea abajo) se observa que hablante 1 dice A y después el hablante 2 dice B; no es necesario que el hablante 2 dice otra cosa o aún el contrario de hablante 1, sino 2 podría sencillamente confirmar una oración de 1. Por ejemplo, 1: "¡Qué cielo azul!" - 2: "¡Qué cielo azul!".

A continuación mostramos la lista de todas las hipótesis actualmente disponible en AQUAD que se activan al seleccionar la opción "*Utilizar vínculos prefabricados*" -> "*Comparar 2 hablantes*" en el menú "*Vínculos*":

Seleccionar estructura de vínculo (S1 = hablante 1; S2 = hablante 2) :

- ☐ S1: Cód. A -> S2: Cód. B
- ☐ S1: Cód. A -> S2: Cód. B y Cód. C
- ☐ S1: Cód. A -> S2: Cód. B o Cód. C
- ☐ S1: Cód. A y Cód. B -> S2: Cód. C
- ☐ S1: Cód. A y Cód. B -> S2: Cód. C y Cód. D
- ☐ S1: Cód. A y Cód. B -> S2: Cód. C o Cód. D
- ☐ S1: Cód. A o Cód. B -> S2: Cód. C
- ☐ S1: Cód. A o Cód. B -> S2: Cód. C y Cód. D
- ☐ S1: Cód. A o Cód. B -> S2: Cód. C o Cód. D
- ☐ S1: Cód. A y (Cód. B o Cód. C) -> S2: Cód. D
- ☐ S1: Cód. A y (Cód. B o Cód. C) -> S2: Cód. D y Cód. E
- ☐ S1: Cód. A y (Cód. B o Cód. C) -> S2: Cód. D o Cód. E
- ☐ S1: Cód. A y (Cód. B o Cód. C) -> S2: Cód. D y (Cód. E o Cód. F)

En AQUAD, se ofrecen 13 "enlaces" hipotéticos para comprobar las conexiones entre los enunciados de dos hablantes. Se insertan códigos concretos en las especificaciones abstractas. A continuación se muestra una lista de todas las hipótesis disponibles actualmente en AQUAD.

1. Hablante 1 dice A, luego hablante 2 dice B en su próximo segmento de datos. ¿Verdadero para qué casos?
2. Hablante 1 dice A, luego hablante 2 dice B y C en su próximo segmento de datos. ¿Verdadero para qué casos? (Se podría contestar por ejemplo la suposición: "Si hablante 1 dice algo sobre X, el hablante 2 se refiere también a X y describe un ejemplo concreto.").
3. Hablante 1 dice A, luego hablante 2 dice B o C (O inclusivo) en su próximo segmento de datos. ¿Verdadero para qué casos?

4. Hablante 1 dice A y B, luego hablante 2 dice C en su próximo segmento de datos. ¿Verdadero para qué casos?
5. Hablante 1 dice A y B, luego hablante 2 dice C y D en su próximo segmento de datos. ¿Verdadero para qué casos?
6. Hablante 1 dice A y B, luego hablante 2 dice C o D (O inclusivo) en su próximo segmento de datos. ¿Verdadero para qué casos?
7. Hablante 1 dice A o B (O inclusivo), luego hablante 2 dice C en su próximo segmento de datos. ¿Verdadero para qué casos?
8. Hablante 1 dice A o B (O inclusivo), luego hablante 2 dice C y D en su próximo segmento de datos. ¿Verdadero para qué casos?
9. Hablante 1 dice A o B (O inclusivo), luego hablante 2 dice C o D (O inclusivo) en su próximo segmento de datos. ¿Verdadero para qué casos?
10. Hablante 1 dice A y (B o C), luego hablante 2 dice D en su próximo segmento de datos. ¿Verdadero para qué casos? (Por ejemplo: "Hablante 1 cuenta de su fin de semana a la playa y describe el viento fuerte y/o las olas enormes, luego hablante 2 aporta una descripción de las olas").
11. Hablante 1 dice A y (B o C), luego hablante 2 dice D y E en su próximo segmento de datos. ¿Verdadero para qué casos?
12. Hablante 1 dice A y (B o C), luego hablante 2 dice D o E en su próximo segmento de datos. ¿Verdadero para qué casos?
13. Hablante 1 dice A y (B o C), luego hablante 2 dice D y (E o F) en su próximo segmento de datos. ¿Verdadero para qué casos?

Si quiere trabajar con cualquiera de estas hipótesis, solamente tiene que seleccionarla activando el botón de opción que la precede en la ventana correspondiente. Vamos a ver un ejemplo de como se completaría la hipótesis tercera. Al hacer clic sobre el botón de selección de esta hipótesis, se abre la ventana que mostramos a continuación.

Esta permite seleccionar dos códigos de hablante (un del hablante 1, el otro del hablante 2) y un código conceptual para hablante 1, dos códigos conceptuales vinculados por el "O lógico de inclusión" para hablante 2. Deseamos saber si se produce un vínculo entre las preguntas del entrevistador y las respuestas del profesor en los archivos del proyecto "entrevistas". Es que suponer que los profesores hablen sobre problemas o algún dilema, p.ej.

respecto a la infraestructura de su centro escolar y/o de su clase/sus alumnos en general cuando el entrevistador quiere saber algo que causa problemas (código "*problemas*").

Podemos ver arriba la forma como escribiríamos nuestra hipótesis. En primer lugar haríamos clic sobre el primer código de hablante en la lista de la izquierda, este código se colocaría en la primera casilla; a continuación haríamos lo propio con el segundo hablante y con los códigos conceptuales que determinan nuestra suposición. Pulsamos el botón "*Continuar*" y AQUAD proporcionará los resultados.

7.3.2 Cómo construir vínculos

Como se ha visto anteriormente, es posible establecer vínculos con elementos de comparación fijos. Así, existen "fórmulas de vinculación" en las que uno inserta sus propios códigos como "variables". Todas estas variables están vinculadas con la lógica "Y" u "O". En general, esto es suficiente. Sin embargo, si uno llega a hipótesis según las cuales, por ejemplo, un hablante se refiere a un tema, debe construir sus propias hipótesis de enlace y, en ocasiones, utilizar también "NO" como operador lógico. A continuación, Leo Gürtler ofrece ejemplos y sugerencias a partir de sus experiencias con los análisis cualitativos en el contexto de su disertación sobre cómo se pueden crear de forma sensata las propias hipótesis de enlace. Para ello, se combinan las posibilidades de la opción del programa "construir vínculos" y "códigos de secuencia" (insertar en archivos de códigos).

Antes de adentrarnos en estas vinculaciones complejas, echaremos un vistazo al sencillo ejemplo de la vinculación de los códigos "problemas" y "DP sí-mismo" (véase el ejemplo anterior) para ver cómo uno mismo puede construir la hipótesis de su vinculación y qué resulta al comprobarla:

Operador	Código
Y	problemas
Y	DP sí-mismo

clase - alumnos
 confianza-reglas
 diferenciación
 dilema
 directivas-libertad
 DP aprender/prof.
 DP exp. doc. previas
 DP opiniones
 DP sí-mismo
 ejemplo
 expectativas
 experiencia

Distancia máxima entre segmentos:
 Tamaño de filas: 3

Código de secuencia:

Esta construcción puede guardarse para su posterior reutilización haciendo clic en "*Guardar*". El botón "*Vamos*" inicia el análisis y muestra el resultado.

Ahora, el uso de las construcciones de hipótesis propias en un contexto complejo: En su investigación sobre las teorías subjetivas del humor (respuestas codificadas de los estudiantes en un cuestionario; Gürtler, 2004), surgió la conjetura de que existe una relación entre la respuesta a dos preguntas:

Pregunta 7 - "¿Crees que hay suficiente humor en clase?" y

Pregunta 8- *"Si usted mismo pudiera cambiar su enseñanza, ¿qué haría para que fuera más humorística?"*

Era interesante averiguar si los estudiantes que respondieron insatisfechos a la pregunta 7, es decir, que no estaban satisfechos con la calidad o la cantidad de humor en el aula, podrían no tener ninguna sugerencia de mejora que ofrecer a la pregunta 8. Esto indicaría que la insatisfacción puede darse con una tendencia a la pasividad en los patrones de respuesta y que habría que distinguir las respuestas de aquellos estudiantes que también expresan insatisfacción en la pregunta 7, pero que sí ofrecen sugerencias de mejora en la pregunta 8 (por ejemplo, cambios en el clima del aula, el estilo de enseñanza de los profesores o los métodos utilizados). Para separar e identificar cualitativamente estos grupos, que podrían ser subelementos de un tipo, se establecieron los siguientes códigos de secuencia, a los que se dio un nombre único y se insertaron en los archivos de códigos:

SeqCode 1: "F7_Insatisfacción"

Y	Código de hablante:	/\$Pregunta7
Y	Código:	StatusQuo: no hay suficiente humor
NO	Código:	StatusQuo: basta de humor

SeqCode 2: "F7_Satisfacción"

Y	Código de hablante:	/\$Pregunta7
Y	Código:	StatusQuo: basta de humor
N	Código:	StatusQuo: no hay suficiente humor

SeqCode 3: "F8_ninguna_respuesta"

Y	Código de hablante:	/\$Pregunta8
Y	Código:	Datos que faltan
O	Código:	no sé/no me importa

SeqCode 4: "F8_Sugerencias"

Y	Código de hablante:	/\$Pregunta8
Y	Código:	Cambiar las estructuras (enseñanza/métodos)
OR	Código :	Promover el clima/las relaciones
OR	Código :	Cambios al nivel institucional

Como puede ver, cada uno de estos códigos de secuencia está claramente vinculado a una de las preguntas (había nueve en total) del cuestionario. Todos los resultados positivos de la búsqueda posterior arrojaron resultados, que por lo tanto también están relacionados exclusivamente con la respuesta a esta pregunta. En un segundo paso, se establecieron otros códigos de secuencia. Estos se veían así:

SeqCode 5: "Vgl_F7<->F8_inconsistent1"

Y	Código:	F7_Insatisfacción
Y	Código:	F8_sin respuesta

SeqCode 6: "Vgl_F7<->F8_inconsistent2"

Y	Código:	F7_Insatisfacción
NO	Código:	F8_sugerencias

SeqCode 7: "Vgl_F7<->F8_inconsistent3"

Y	Código:	F7_satisfacción
---	---------	-----------------

Y Código: F8_sugerencias

SeqCode 8: "Vgl_F7<->F8_consistent1"

Y Código: F7_Insatisfacción

Y Código: F8_sugerencias

Este esquema de codificación permite distinguir a los estudiantes que

- están insatisfechos con la situación de la escuela, pero no tienen ninguna respuesta (= "datos perdidos") a la pregunta sobre las posibilidades de mejora. Esto es incoherente y podría dar pistas sobre el potencial de acción experimentado en el contexto de la escuela o la voluntad o capacidad de expresar o aportar sus propias ideas y sugerencias.
- están insatisfechos, pero no presentan ninguna sugerencia. Esto también es inconsistente como el ejemplo anterior, pero aquí aparentemente se dan otras expresiones y por lo tanto una respuesta en absoluto, lo que es cualitativamente diferente de SeqCode 5.
- están satisfechos con la situación de la escuela, lo que podría significar que no quieren cambiar nada, pero aún así presentan sugerencias de mejora (aquí la pregunta es: ¿estos estudiantes no están realmente satisfechos o son más bien creativos y comprometidos, lo que habría que investigar en un análisis separado).
- están insatisfechos pero tienen sugerencias de mejora que ofrecer. Esto es coherente en el sentido de que algo se percibe como incongruente, pero también se hacen al menos esfuerzos cognitivos para abordar esta incongruencia.

La combinación de la codificación de secuencias con la ya existente es excelente para responder a preguntas más complicadas sin ambigüedades. El ejemplo debe enfatizar además el uso inteligente y planificado de los códigos de los hablantes. Los códigos de los hablantes tienen la útil propiedad de separar claramente unas partes del texto de otras, lo que desempeña un papel crucial en la comparación de hipótesis cualitativas en la formación de tipos. Además de los hablantes reales, puede utilizar códigos de hablante para las preguntas de un cuestionario (véase el ejemplo) o diferentes áreas como las emociones, las cogniciones, las acciones. Este procedimiento es un excelente punto de partida para la formación de tipos, así como una base para un posterior análisis de implicantes (cf. Ragin, 1987). Especialmente en el análisis de implicantes (véase más adelante; QCA: Qualitative Comparative Analysis), se beneficia no sólo de relacionar los códigos de contenido individuales entre sí, sino de utilizar proposiciones (por ejemplo, sujeto - predicado - objeto, es decir, parte/acciones) que se han operacionalizado mediante la codificación de la secuencia.

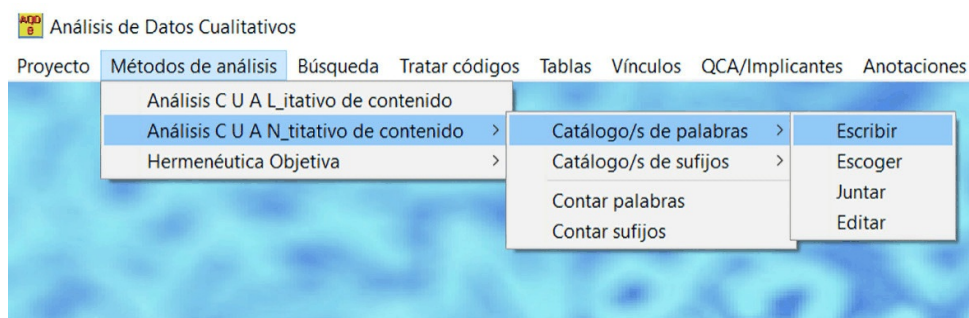
Capítulo 8: Análisis de contenido cuantitativo

8.1 Recuento de palabras

No podemos entrar aquí en la polémica sobre el uso de enfoques cuantitativos en el análisis de datos cualitativos. Sin embargo, no cabe duda de que complementar el análisis cualitativo con el cuantitativo puede ser útil para algunas cuestiones de investigación. Si se han identificado las palabras clave como indicadores de determinadas afirmaciones en los textos, entonces, al contar estas palabras y comparar sus frecuencias en los textos, se pueden obtener, al menos, indicaciones iniciales de puntos focales de significado. Otros analistas de textos irían más allá y utilizarían los resultados de un análisis de frecuencias a nivel de palabras como datos iniciales de los procedimientos estadísticos pertinentes,

como los análisis de cluster. En Vorderer y Groeben (1987) ya se puede encontrar un análisis detallado de las frecuencias de las palabras.

El acceso a las funciones de recuento en textos se encuentra en el menú principal de AQUAD en el grupo de funciones "*Métodos de análisis*" -> "*Análisis CUANtitativo de contenido*". La siguiente figura ofrece una visión general de cómo se pueden crear listas de palabras críticas. El "*Contar palabras*" se puede seleccionar en el submenú anterior de las subfunciones.



Antes de empezar a contar palabras, debe introducir las palabras que desea contar o introducir el nombre de una lista de palabras correspondiente que ya esté disponible. En este caso, el programa no busca las palabras como "partes del texto", sino que desglosa todo el texto en palabras individuales. Todos ellos se convierten en minúsculas. Esto significa que el ordenador sólo cuenta las palabras enteras. Si quiere contar sólo determinadas palabras con listas de palabras, el programa sólo tiene en cuenta las palabras que están contenidas en la lista. Por ejemplo, si has introducido la palabra "amigo", el ordenador no cuenta "amistad" o "círculo de amigos", etc.

Como segunda opción, puede hacer que el programa liste todas las palabras de un texto concreto (o de varios textos). A partir de la lista de palabras individuales ordenada alfabéticamente en la pantalla, se seleccionan, haciendo clic sobre ellas, las que se van a recopilar en una lista de palabras para realizar análisis de frecuencia limitada. Es aconsejable crear primero una lista de palabras del texto para cada texto crítico por la función "*Escoger*".

Si ha guardado palabras de varios textos en diferentes listas de este modo, puede "*Juntar*" exactamente estas (o algunas) listas de palabras con la tercera subfunción.

La última opción es "*Editar*" las listas de palabras. Para ello, se carga en la pantalla una lista ya disponible y se borran las palabras que no se quieren incluir en un posterior análisis de frecuencias. Por supuesto, también puedes añadir más palabras.

Para su posterior reutilización, todas las listas se guardan con un nombre de libre elección. AQUAD añade la extensión ".cwo" (catálogo de palabras) al nombre del archivo para su reconocimiento automático. Las listas de palabras en forma de campos de palabras o "léxicos de significado" ayudan a obtener rápidamente una visión general de las áreas de contenido relevantes en los textos de un proyecto.

Si finalmente hacemos clic en "*Contar palabras*", aparece primero la siguiente ventana, en la que se definen las condiciones del análisis y la salida de resultados. En concreto, AQUAD ofrece las siguientes posibilidades:

Contar palabras (Proyecto: Entrevista)

Escoger Condiciones:

☐ separado por hablantes

☐ Lista (P/S) -> buscar

☐ Lista (P/S) -> excluir

Mostrar cada resultado

☐ Mostrar cada resultado

Mostrar resultados, si

☐ Frecuencia f > 0

☐ Frecuencia f < 2

Continuar

Cancelar

Ayuda

Si los textos están subdivididos con códigos de hablante, el recuento de palabras puede hacerse, por supuesto, por separado por hablante. (Nos gustaría recordar aquí que "hablante" puede interpretarse de forma creativa; por ejemplo, al analizar las respuestas abiertas en un cuestionario, se podrían marcar las preguntas individuales con códigos de hablante y luego evaluarlas por separado).

Sin embargo, los oradores individuales, por ejemplo el moderador / la moderadora de una discusión de grupo, ya pueden ser excluidos del análisis cuantitativo durante la codificación mediante el código de control "\$no contar".

Si se han creado listas de palabras, estos archivos se muestran para su selección a la derecha en la ventana, que aquí sigue vacía, si se marca una de las casillas "Lista (P/S) -> buscar" o "Lista (P/S) -> excluir". De esta manera, usted decide si sólo se cuentan las palabras de la lista seleccionada o se excluyen exactamente estas palabras del recuento. Sin seleccionar una lista, se cuentan todas las palabras de los textos.

Para la visualización/salida de los resultados, vemos tres casillas de verificación en la parte inferior izquierda: Podemos mostrar todos los resultados - entonces se omiten las siguientes opciones. O bien seleccionamos un límite inferior ($f > 0$) y/o un límite superior ($f < 2$) para la salida de los hallazgos, donde el "0" y el "2" pueden cambiarse según sea necesario.

El recuento se realiza siempre en todos los archivos de un proyecto. aquí en las cuatro transcripciones de entrevistas del proyecto "entrevista". Si sólo quiere tener en cuenta determinados archivos, puede crear una definición de proyecto en la que sólo aparezcan estos archivos.

¿Cómo es el resultado si dejamos que todas las palabras cuenten? El primer resultado que obtenemos es una lista ordenada alfabéticamente de todas las palabras. Delante de cada palabra se introduce su frecuencia en el texto correspondiente. A la derecha se encuentra la misma lista, pero ordenado por frecuencias de las palabras. Aquí un pequeño extracto del principio de la lista:

```

work))) - Editor
Datei Bearbeiten Format Ansicht Hilfe
| | | entre_1.txt
1: ----
12: -
1: 100
14: a
1: abro
1: abuela
1: acabo
1: acostumbrada
1: adaptarse
3: además
1: adquiere
1: ahí
1: ahora
1: aislándola
1: al
1: algo
1: algunos
1: allí
1: alta
10: alumnos
1: ----
1: 100
1: abro
1: abuela
1: acabo
1: acostumbrada
1: adaptarse
1: adquiere
1: ahí
1: ahora
1: aislándola
1: al
1: algo
1: algunos
1: allí
1: alta
1: aquel
1: asignaturas
1: aspecto
1: aula

```

A continuación se presenta una tabla con el número total de palabras de cada texto, el número de palabras diferentes por texto y la redundancia del texto, es decir, la relación entre el número total de palabras y el número de palabras diferentes:

Contar palabras (Proyecto: Entrevista)

Archivos del proyecto:

- entre_1.txt
- entre_2.txt
- entre_3.txt
- entre_4.txt

Archivo de texto	Total	Palabras	Redundancia
entre_1.txt	858	313	0.64
entre_2.txt	1065	365	0.66
entre_3.txt	964	384	0.60
entre_4.txt	894	356	0.60

Atención:

Tabla de formato CSV:
¿Quiere guardar esta tabla?

SI NO

OK Cancelar

Nota: Si quiere hacer cálculos estadísticos con las frecuencias, debe corregir las frecuencias de las palabras críticas en función del número total de palabras de un texto. Puede haber una gran diferencia en las conclusiones si, por ejemplo, la palabra "yo" aparece doce veces en un texto que consta de un total de 700 palabras o de 1400.

Como ejemplo de la necesidad de corregir los datos de frecuencia, veamos las frecuencias de las palabras contenidas en la lista de palabras "alumnos.cwo" al contar con las cuatro transcripciones de las entrevistas. En esta lista de palabras, las palabras que los profesores utilizaron para referirse a sus alumnos se recopilaron a partir de las cuatro entrevistas; analizamos la columna "Total":

Contar palabras (Proyecto: Entrevista)

Archivos del proyecto:

- entre_1.txt
- entre_2.txt
- entre_3.txt
- entre_4.txt

Archivo de texto	alumno	alumnos	chaval	chavales	niño	niños	Total
entre_1.txt	0	10	0	2	0	0	12
entre_2.txt	2	6	0	1	0	6	15
entre_3.txt	0	1	0	0	1	6	8
entre_4.txt	0	1	2	0	4	5	12

Atención:

Tabla de formato CSV:
¿Quiere guardar esta tabla?

SI NO

OK Cancelar

"alumnos" y terminos relacionados aparecieron 12 veces en la entrevista 1 (858 palabras), 15 veces en la entrevista 2 (1065 palabras), sólo 8 veces en la entrevista 3 (964 palabras) y 12 veces en la entrevista 4 (894 palabras). Tras tener en cuenta la duración de las entrevistas y corregirlas en consecuencia, se confirma la diferencia entre las entrevistas en lo que respecta a la atención a los alumnos. Pero lo que esto puede significar es una pregunta que debe ser respondida cualitativamente...

Al final, no debemos olvidar guardar la lista de frecuencias. Para ello, hacemos clic en el módulo "*Archivo*" de la parte superior de la ventana de resultados y, a continuación, en "*Guardar como*". Esto abre el diálogo habitual de guardado. Seleccione el subdirectorio "...res" del directorio de los archivos AQUAD (la ruta de los resultados) para guardar. Se recomienda utilizar el nombre de la lista de palabras – en este caso "alumno" – como nombre de archivo para guardar las listas de frecuencias y las letras ".frq" (por "frecuencia") como extensión. Esta es la forma más fácil de encontrar los resultados más tarde.

8.2 Contar sufijos

Para cuestiones especiales, como la determinación de los estilos cognitivos a partir de las producciones verbales de las personas (cf. Günther, 1987), parece adecuado utilizar un algoritmo de búsqueda y recuento que sólo registre las terminaciones de las palabras o los sufijos (como en alemán "-heit", "-keit", "-chen" o "-lein"). Una vez más, hay que crear un directorio en el que se introducen todos los sufijos que se quieren contar. A continuación, seleccione "*Contar sufijos*" y obtendrá las frecuencias. Al introducir los sufijos, debes pensar también en las formas plurales.

8.3 Cómo excluir partes del texto del análisis de palabras

Más arriba mencionamos cómo se pueden excluir partes de un texto del análisis con un código de control. A estas alturas debería haber quedado claro por qué esta opción es a menudo urgente. En los análisis de entrevistas, por ejemplo, se mostrarían o contarían todas las expresiones verbales del entrevistador.

Aquí mostramos un pequeño ejemplo de texto de una interacción profesor-alumno en el que se iba a analizar la frecuencia con la que los alumnos se refieren a los objetos críticos del problema, las esferas. Por lo tanto, hemos ocultado los enunciados del profesor con el código de control "\$no contar". De este modo, se excluyen del análisis las expresiones del profesor, pero no las de los alumnos:

54 ...	
55 PROF. Cuántas balas más podemos ponerle,	\$no contar 55 - 56
56 ¿Que hay sentido?	
57 ALUMNO ¿Eh?	
58 PROF. Cuántas bolas más podemos comparar	\$no contar 58 - 59
59 ¿Cuántas bolas hay en cada lado? ...	
60 ALUMNO Dos - tres - y cuatro, o una bala ...	
61 PROF. ¿Y cuál es la más eficiente?	\$no contar 61 - 61
62 ALUMNO Cuatro	
63 PROF. ¿Por qué?	\$no contar 63 - 63
64 ...	

Capítulo 9: Trabajar con memos

9.1 ¿Para qué sirven los memos?

Todas las ideas, preocupaciones, posibles contradicciones, etc. que surjan durante el trabajo, pero también otras notas del proyecto, referencias cruzadas, referencias bibliográficas que no quieras olvidar, pueden registrarse directamente en AQUAD y volver a buscarse específicamente. Te recomendamos encarecidamente que utilices ampliamente la función de notas de AQUAD. Debido a la gran importancia de los memos en los procesos de comparación constante de textos y de análisis de construcción de teorías, también se han incorporado a AQUAD algoritmos especiales de búsqueda para ayudar a los investigadores a reconectar posteriormente los memos con los textos, segmentos, codificaciones, etc. que dieron lugar a la idea que anotaron.

Basándonos libremente en Miles y Huberman (1991, p. 69), destacamos que la recopilación y el análisis de textos es tan emocionante y la codificación a menudo tan agotadora que es demasiado fácil perderse en la abundancia de detalles interesantes, como son "... las frases mordaces, la intrigante personalidad de un informante clave, la reveladora viñeta en el tablón de anuncios del pasillo, las conversaciones informales después de una reunión importante. Se olvidan de pensar, de buscar el significado más profundo y general de lo que ocurre, de explicarlo de forma conceptualmente coherente. ...". Por lo tanto, es casi indispensable anotar todo lo que pasa por la mente sobre las categorías centrales y sus posibles conexiones durante las entrevistas, las observaciones, las conversaciones, etc. y luego al codificar.

Hace unos años, Glaser (véase Glaser 1992, p. 108 y ss.) y Strauss (véase Strauss & Corbin 1990, p. 197 y ss.) tuvieron una discusión sobre la importancia de los memos en el proceso de los análisis cualitativos. Coinciden en que los memos deben acompañar el proceso de investigación desde los primeros inicios hasta el informe final, y que no hay que abstenerse de anotar las ideas inmediatamente en cualquier fase. Los memos son soportes importantes para la reflexión teórica. El abandono de los memos es evidente en el producto final del proyecto: "una teoría cuyos conceptos carecen de densidad y/o sólo están vagamente relacionados" (Strauss & Corbin, 1990, p. 199).

Sin embargo, Strauss y Corbin se dedican a canonizar los memos espontáneos, creativos, quizás al principio salvajemente especulativos; empiezan a distinguir entre memos, notas de codificación, notas teóricas, notas operativas, diagramas y diagramas lógicos, y desarrollan procedimientos específicos para cada uno. En este punto, recomendamos seguir la objeción de Glaser (1992, p. 109) de que uno no puede desarrollar por adelantado y en general un sistema de si y por qué hay que prestar atención a características muy específicas en cualquier teoría que uno esté construyendo - hasta que uno se encuentra con estas mismas características. Por otra parte, estos sistemas de rasgos pueden tener funciones heurísticas útiles (cf. Huber 1992, p. 145 y ss.), pero de ninguna manera garantizan o fuerzan la construcción de un edificio teórico particular como un esquema descriptivo o exhaustivo. O dicho de otro modo: según el enfoque de generación de una "teoría fundamentada", lo ideal sería que "... el analista comenzara sin ningún tipo de suposición estructural y dejara que los conceptos se estructuraran por sí mismos en el proceso de su emergencia" (Glaser 1992, p. 110); pero si ahora no quiere emerger nada, algunas sugerencias sobre los aspectos bajo los que se podría seguir mirando un texto pueden ser bastante útiles.

Como conclusión, sostenemos: anote todas sus ideas en todas las fases del análisis cualitativo como memos. No intente seguir un sistema rígido de notación desarrollado independientemente de sus textos.

9.2 Cómo escribir memos en AQUAD

AQUAD ofrece dos formas posibles de acceder a sus funciones de memo, dependiendo de los módulos con los que estés trabajando:

1. Hay una opción "Memos" en el menú principal, y
2. hay una columna "Memo" en la parte izquierda de la tabla de codificación (en las opciones de "Análisis CUAN_titativo de contenido" del menú principal).

Especialmente al codificar, a menudo querrás anotar una idea. Por lo tanto, tratemos esto primero.

9.2.1 Escribir notas durante se codifica

Al codificar textos, imágenes, grabaciones de audio o vídeo, creas tus propios archivos de código para tus archivos de datos. Al mismo tiempo, puedes activar la función de memorándum siempre que sea necesario haciendo clic en la columna de memos ("M") de la parte izquierda de la ventana de codificación.

En las líneas de la tabla de codificación en las que aún no se ha adjuntado ninguna nota, la columna "Memo" está vacía, en las demás líneas se introduce una "X" en esta columna. Al hacer clic en este signo se abre la ventana de memos con la primera de las posibles notas asociadas a esta fila.

Al hacer clic en una celda vacía de la columna "Memo" se abre una ventana vacía en la que se puede escribir una anotación:

En esta casilla, algunas entradas ya se realizan automáticamente: (1) el nombre del archivo al que está vinculado este memo, (2) el código vinculado a estas filas y (3) el número de la fila en la que se hizo clic en la celda de la columna "Memo". Los usuarios han solicitado poder cambiar las dos primeras entradas (doble clic en las ventanas de texto). La modificación de estos criterios se ve facilitada por el hecho de que al hacer doble clic en los campos de

nombre de archivo y código se abre una ventana adicional en cada caso (directorio de archivos del proyecto; registro de código), en la que se selecciona el contenido adecuado haciendo clic. Normalmente no se le ocurriría cambiar el nombre del archivo al que se refiere su anotación, pero si fuera necesario... Para volver a cerrar estas ventanas, basta con hacer doble clic. Según la lógica de "*escribir memos durante se codifica*", la posición/número de línea (como determinantes de un segmento de archivo) sólo debe cambiarse durante la codificación haciendo clic en la columna "Memo" en otra línea.

En la línea de entrada bajo el título amarillo "*Introducir texto y buscar en los memos*" puedes introducir una parte del texto o sólo una palabra clave si luego quieres encontrar aquellos memos que contengan esta parte del texto o esta palabra clave al "navegar" por los memos.

¿Y el texto de la nota? Tienes una ventana de entrada amarillenta delante de ti, pero la longitud real de una anotación es prácticamente ilimitada. Por ejemplo, puedes copiar un archivo de texto completo en el cuadro de notas, si es que eso tiene sentido. En otras palabras, no hay restricciones a la hora de escribir memos. La copia de pasajes de texto y de partes de texto de otros memos se realiza haciendo clic con el botón derecho del ratón en cualquier parte del texto. El texto completo aparece entonces en el editor, donde se puede marcar el pasaje deseado y copiarlo con las funciones de edición. En la ventana de memos, se hace clic con el botón derecho del ratón en el cuadro de notas y se selecciona la función "Pegar". Seguramente también verás aquí la opción de imprimir tus memos.

iiiNo olvides guardar la anotación con "*Guardar*" cuando hayas terminado con tu entrada!!!

9.2.2 Escribir memos desde el menú principal

Las explicaciones del apartado 9.2.1 se aplican a todas las fases del análisis. Una diferencia esencial, sin embargo, es que desde el menú principal no se da ninguna entrada (por ejemplo, el nombre del archivo) en ninguna de las ventanas de entrada de un memo, porque en este caso AQUAD no sabe a qué archivo o a qué segmento de datos se va a asignar el memo.

Puede seleccionar las entradas que faltan en las listas que aparecen al hacer clic en la columna de entrada correspondiente (por ejemplo, "*Código*"). Sin embargo, debe introducir directamente en qué línea comienza el segmento de texto al que se asigna la nueva nota.

9.3 ¿Cómo puedo volver a encontrar mis memos?

Esta es una cuestión que hemos tratado a fondo en la programación de AQUAD. A menudo ocurre con las ideas fugaces, que se registran ventajosamente en memos, que uno recuerda una idea tremendamente importante en este o aquel archivo, en relación con un determinado segmento, etc., pero simplemente no recuerda exactamente cuál era la importancia del pensamiento.

Si pudieras recordar los detalles, habrías encontrado el memorándum enseguida. Ahora, por supuesto, puedes leer todos los memos uno por uno - eso también es posible en AQUAD. Sin embargo, el apoyo a una consulta definida es muy útil en esta situación. Si inicias una consulta con el objetivo de recuperar tus memos, todas las entradas visibles en tu ventana de memos pueden servir como criterios de búsqueda.

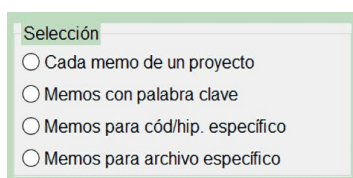
9.3.1 Buscar memos durante se codifica

Mientras esté ocupado codificando un archivo concreto, una parte considerable de su trabajo consistirá en la "*comparación permanente*" (Glaser y Strauss, 1967). Cuando se atribuye un determinado significado a un segmento de archivo y se codifica, a menudo se recuerda otro segmento que parece muy similar o muy diferente. ¿Cómo codificaste ese segmento, y por qué justo así? Así que volverás a ese segmento y observarás de nuevo con detenimiento la(s) codificación(es) asociada(s) a él. O puede recordar que ya ha utilizado un código concreto en una situación similar. ¡Las notas con sus reflexiones al codificar este segmento o utilizar un código específico serían muy útiles ahora! Busque el segmento de datos o el lugar donde ha utilizado el código crítico, haga clic en la celda correspondiente de la tabla de codificación en la columna "Memo" y lea el memo que se refiere al segmento de datos correspondiente y/o al código.

Si haces clic con el botón derecho del ratón en cualquier lugar del campo de texto de la anotación, el contenido se transfiere al editor y puede guardarse externamente desde allí, imprimirse, copiarse y pegarse en otro lugar.

Hojea las notas

La opción "*Hojea*" requiere una explicación algo más detallada. Si busca memos específicamente según ciertos criterios mientras codifica, utilice este botón. El segundo clic hace que aparezca la lista de posibles criterios de búsqueda, pero en tres casos aquí debe introducir primero este criterio en el formulario de memos antes de pulsar el botón correspondiente, p.ej. se tiene que entrar el nombre del archivo antes de hacer clic sobre el botón "*Memos para archivo específico*".



Selección

- ☐ Cada memo de un proyecto
- ☐ Memos con palabra clave
- ☐ Memos para cód/hip. específico
- ☐ Memos para archivo específico

Si quiere ver todos los memos del proyecto, sólo tiene que hacer clic en la primera opción, pero entonces obtendrá una lista de todos los memos, ordenados por archivos y números de línea.

Para la segunda opción, antes de desplazarse a la línea bajo el título amarillo "Introducir texto y buscar en los memos", debe introducir una palabra clave o parte de una frase que sirva de criterio para mostrar aquellos memos (de todo el proyecto) en los que el criterio esté contenido. Incluso al redactar los memos, puede utilizar encabezados específicos (por ejemplo, "definición de código") para asegurarse de encontrar las entradas específicas todas juntas.

Para la tercera opción, antes de desplazarse, seleccione el código deseado (haga doble clic en la línea de código) para el que ha escrito memos que ahora quiere volver a encontrar.

Para la cuarta opción, antes de navegar, seleccione el nombre del archivo (haga doble clic en la línea del archivo) para el que ha escrito memos que ahora quiere recuperar.

Los resultados de la búsqueda se muestran en el editor. Si desea guardar o imprimir los memos encontrados externamente, encontrará los comandos correspondientes en la barra de herramientas situada en la parte superior de la ventana de salida del editor, en "*Archivo*". Puedes pegar las partes seleccionadas o todos los memos seleccionados en otros memos o textos abiertos en el editor en "*Editar*" (luego "*Copiar*" y "*Pegar*").

9.3.2 Búsqueda de memos desde el menú principal

Si activa la opción "Anotaciones" del menú principal, tendrá acceso inmediato a todos los memos asociados a su proyecto actual. Se aplican las mismas reglas que para la búsqueda de notas durante la codificación (véase más arriba).

Capítulo 10: Análisis de implicantes (Análisis comparativo cualitativo)

Además de la posibilidad de comprobar las hipótesis sobre las relaciones de significado, los procedimientos para la minimización lógica de las configuraciones condicionales representan otra característica especial de AQUAD. Las ideas básicas de los meta-análisis cualitativos o las comparaciones de casos individuales sistematizadas con él se basan en el trabajo de Charles Ragin (Ragin, 1987).

En resumen, sólo cabe repetir aquí que con el componente "QCA/implicantes" las configuraciones de significado específicas de cada caso se transforman en valores de verdad de una lógica binaria (característica se aplica/no se aplica) y se vinculan entre sí según las reglas del álgebra booleana. Una característica sirve como criterio de comparación. La comparación de todos los casos en los que este criterio es "verdadero" (o "falso") conduce a la reducción de las configuraciones difíciles de entender a los principales implicados de este criterio. Dado que a menudo hay más de un patrón de implicación de este tipo en una investigación, la redundancia lógica puede reducirse aún más analizando los implicantes esenciales y delimitando los implicantes secundarios, lo que a veces es necesario en estructuras condicionales complejas. El resultado representa generalmente configuraciones significativas de condiciones. A continuación se describen los pasos y las posibilidades al utilizar el módulo "QCA/Implicantes".

10.1 Escribir una tabla de datos

AQUAD permite realizar directamente minimizaciones lógicas de las configuraciones de condiciones de múltiples estudios o casos, utilizando datos cuantitativos y/o cualitativos. Se pueden combinar estos datos en una tabla, es decir, escribir, editar tablas existentes e imprimir tablas de datos. Los datos cualitativos, por ejemplo, "*Significado A fuertemente expresado*" (es decir, "*verdadero*"), pueden introducirse directamente en la tabla de valores de verdad. Sin embargo, podría ser ventajoso representar los datos cualitativos en forma de números naturales y luego introducir, por ejemplo, el valor numérico "9" para "*verdadero*" y "1" para "*falso*".

En consecuencia, se escribiría la interpretación "*significado A débil / no pronunciado en absoluto*" como "1" o como "*falso*" en la tabla de datos. AQUAD contempla ambas posibilidades.

Aquí demostramos cómo introducir tanto datos cuantitativos como información cualitativa expresada en valores numéricos. Para ilustrar los pasos, presentamos un ejemplo ficticio de meta-análisis de la relación entre el rendimiento escolar y el tamaño de las clases. Suponemos que se han encontrado 12 estudios relevantes en la literatura. Supongamos además que se registraron cuatro condiciones en los 12 estudios: A = tamaño de la clase, B = aptitud, C = duración del experimento y D = rendimiento escolar. Se supone que nuestros datos contienen tanto valores cualitativos (A: clase evaluada como grande o pequeña) como medidas cuantitativas (B: nivel de aptitud de la clase, C: duración del ensayo, D: rendimiento). Los datos de los casos individuales darán lugar posteriormente a siete configuraciones diferentes de condiciones. No tenemos que preocuparnos por distinguir estos patrones de configuración. Sólo introducimos los resultados de todos los casos individuales sin clasificarlos nosotros mismos y

dejamos que el ordenador los reduzca a diferentes configuraciones de condiciones. Estas son las diferentes condiciones posibles (clases de datos) en este estudio:

A: Tamaño de la clase - "pequeña" o "grande"- según la información de los informes sobre el estudio.

9 = clase pequeña ("verdadera" en el sentido de la hipótesis de los estudios) / 1 = clase grande ("falsa").

B: Resultados (valores estándar) de una prueba de capacidad (media aritmética de los valores individuales de una clase)

C: Duración del estudio (duración en semanas)

D: Resultados (valores estándar) de una prueba de rendimiento (media aritmética)

Los primeros pasos no necesitan explicación: se selecciona la opción de menú "QCA/Implicantes" y allí el submenú "Escribir tabla de datos".

Las dos flechas situadas en la esquina superior izquierda de la ilustración llaman la atención directamente sobre los dos parámetros de la tabla de datos que tiene que establecer primero:

- Número de condiciones: En la versión actual, el programa puede procesar entre dos y doce condiciones (incluida la condición de los criterios). El parámetro "número de condiciones" define el número de columnas de la tabla. El programa podría considerar más condiciones, pero una ampliación tiene poco sentido dada la complejidad de los resultados previstos y la dificultad de interpretarlos. En nuestro caso, el rendimiento (condición D) sirve de criterio para las comparaciones; las configuraciones de las restantes características A (tamaño de la clase), B (capacidad media) y C (duración de la observación) se prueban como condiciones del criterio D.

En nuestro caso, escribimos el número "4" en el primer campo de entrada o utilizamos la flecha superior para encontrar el "4".

- Número de casos: En este punto tiene que definir cuántos resultados de estudios quiere introducir en la tabla (en este ejemplo: el número de estudios; normalmente introducimos el número de archivos de texto o de interlocutores de la entrevista). En otras palabras, se define el número de filas de la tabla. La única limitación es que se necesitan al menos tres casos para que la comparación sea significativa.

En este ejemplo, introducimos el número 12 porque sólo queremos comparar 12 estudios.

Condiciones 4

Casos 12

	A	B	C	D
1	9	36	32	58
2	1	59	6	62
3	9	65	8	60
4	9	63	20	59
5	9	38	15	39
6	1	32	16	32
7	1	60	24	41
8	1	35	4	37
9	9	37	28	59
10	9	62	22	61
11	9	39	14	38
12	1	60	6	63

OK Cancelar

La tabla se amplía automáticamente en función del número de columnas y filas que defina, por lo que puede empezar inmediatamente a rellenar las celdas con sus datos. Recuerde que en nuestro ejemplo ficticio utilizamos tanto datos cualitativos (condición A) como cuantitativos (condiciones B, C y D). Haga clic en el botón "OK" para finalizar las entradas.

A continuación, se abre una ventana adicional en la que se introduce el nombre con el que se va a guardar la tabla. Cuando haya terminado, pulse "OK" en la pequeña ventana y la tabla se guardará (véase la confirmación).

Encontrará este archivo en dos formatos en el subdirectorio "..\cod" como "*Ejemplo_DT.txt*" en formato de texto y en el subdirectorio "..\res" como "*Ejemplo_DT.csv*" en formato CSV ("_DT" significa Data Table/Tabla de Datos). Ambos ya se han copiado en el directorio que especificó para los archivos de codificación durante la instalación. Si desea experimentar con "*QCA/Implicantes*", encontrará estos archivos en el cuadro de selección correspondiente de la pantalla (véase más abajo).

Un breve comentario sobre el significado de los dígitos de la tabla: Suponemos que no tenemos información más precisa sobre las clases participantes y sólo sabemos si son grandes o pequeñas (condición A). Para la transformación en valores de verdad, introducimos números grandes para las características que se consideran "verdaderas" y números pequeños para las características que se consideran "falsas". Como suponemos que existe una correlación entre las clases pequeñas y el alto rendimiento, asignamos el número 9 a las clases pequeñas (decisión arbitraria) y el número 1 a las clases grandes.

Hay valores métricos disponibles para otras condiciones (B: valores de habilidad; C: duración del experimento; D: valores de rendimiento). Estos valores se introducen tal cual. En un momento posterior, el programa los transforma en valores de verdad.

Se recomienda transformar todas las condiciones en números, incluso si son exclusivamente condiciones cualitativas que se expresaron originalmente con los medios del lenguaje (como "mucho", "poco", "a menudo", etc.). A continuación, introduzca estos números en la tabla (como se ha descrito anteriormente). Este procedimiento ha demostrado ser más cómodo que reducir los resultados a los valores de verdad "verdadero" o "falso" en forma de letras mayúsculas y minúsculas y teclearlos (véase más adelante).

En el ejemplo anterior, la condición A se trata de esta manera. En otros casos podríamos, por ejemplo, escribir el número 9 en la columna correspondiente si tuviéramos que juzgar una afirmación como "Condición... se da" o "Persona ... puede mantener su posición". En el caso contrario, es decir, si la "condición ... no se da" o si la "persona ... no puede hacerse valer", introduciríamos el número 1. Puede elegir casi cualquier valor que desee. Sólo es importante que, al convertir la distribución, se creen valores por defecto que estén por encima o por debajo del porcentaje crítico (corte), que se sustituirán por letras mayúsculas ("verdadero") y minúsculas ("falso") según corresponda.

Por cierto, hay un atajo para pasar de la tabla de frecuencias a la tabla de minimización lógica. Puede llamar a la opción "*Contar códigos*" en el submenú "*Búsqueda*" para producir una lista de frecuencias de aquellos códigos que representan condiciones relevantes para una comparación de los casos de su estudio. Puede guardar esta lista en el editor. Además, la lista se convierte en una tabla de datos y se guarda en formato CSV. El resultado del recuento de frecuencias se guarda bajo un nombre de archivo (..._DT.csv) que Usted introduce para ello. Simplemente seleccione este archivo cuando active la opción "*Convertir tabla de datos -> tabla de verdad*" en el submenú de "*QCA/Implicantes*".

10.2 Convertir datos en valores de verdad

En el camino de una tabla de datos a la tabla de valores de verdad con la que luego se realiza la minimización lógica, la conversión de los datos iniciales en la información reducida "verdadero" o "falso" es el paso decisivo. Para facilitar la interpretación de los distintos resultados del proceso de minimización, AQUAD no trabaja con números binarios ni con los algoritmos correspondientes. En una expresión como "1001" o en el resultado "01*", habría que determinar una y otra vez el significado de cada uno de los valores de verdad (o condiciones eliminadas) en función de la posición en la expresión. En su lugar, AQUAD utiliza las abreviaturas de las letras de las condiciones, es decir, las mayúsculas significan "verdadero" y las minúsculas "falso". En nuestro ejemplo, un caso con la configuración

- A: clase pequeña ("verdadero")
- B: con un nivel de aptitud relativamente alto ("verdadero"),
- c: en la que, tras unas pocas semanas de experimentación ("falsa")
- d: bajo rendimiento ("falso"),

no se representa con la expresión "1100" en la tabla de valores de verdad, sino que se entiende más fácilmente con "ABcd".

Entonces, ¿cómo pasamos de la tabla de datos a las expresiones de valor de verdad como en nuestro ejemplo? Seleccionamos la opción "*Convertir tabla de datos -> tabla de verdad*" en el submenú de "*QCA/Implicantes*" o transformamos nosotros mismos los datos originales en valores de verdad y seleccionamos "*Escribir tabla de valores de verdad*".

En este último caso, el procedimiento es el mismo que cuando se escriben tablas de datos, pero en lugar de números, hay que introducir las expresiones de la condición respectiva (letra en la cabecera de la columna) como mayúsculas ("verdadero") o minúsculas ("falso") en las celdas de la tabla.

Al transformar las tablas de datos, AQUAD realiza la siguiente estrategia de transformación: Para cada condición, los valores de la tabla se estandarizan primero (en todos los casos), es decir, se transforman en valores

predeterminados con $M=100$, $SD=10$. A continuación, el valor de esta tabla intermedia interna se transforma en una letra mayúscula o minúscula según un determinado criterio de corte. El valor del corte se fija en

$$\text{Porcentaje crítica (Corte)} = 50$$

lo que significa que los valores de una condición que están por debajo de 50, es decir, la "mitad inferior", se transforman en el valor de verdad "falso" y se simbolizan con letras minúsculas. En consecuencia, los valores superiores a 50 se reducen al valor de verdad "verdadero" y se simbolizan con letras mayúsculas.

Por supuesto, AQUAD permite cambiar el corte si la pregunta de investigación lo requiere. En la parte inferior de la ventana de transformación, verá un pequeño cuadro que muestra (en caracteres rojos) el valor de corte que se ha establecido. Puedes cambiar este valor paso a paso (en pasos de 5) pulsando sobre él.

Recuerde: Cuanto más alto sea el corte, más datos se asignarán al valor de verdad "falso" (por ejemplo, si la casilla muestra "70", los datos con valores inferiores a 70 se considerarán "falsos" y sólo los valores superiores se aceptarán como "verdaderos". Los cambios en el valor de corte sólo se aplican a la tabla para la que se hizo el cambio. Cuando se introducen datos en una nueva tabla, se utiliza el valor por defecto de "50", hasta que lo ajuste a las necesidades de su estudio.

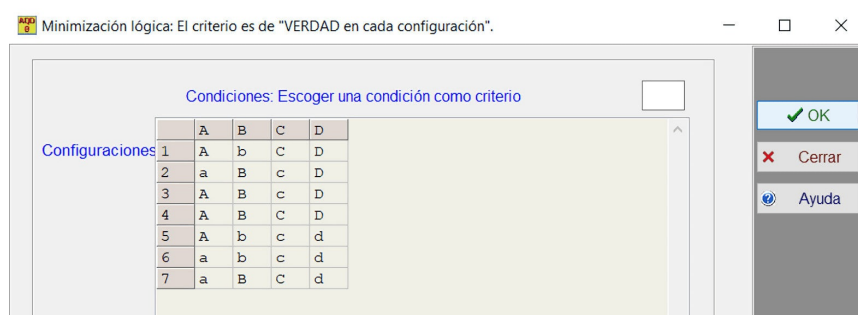
Ahora estamos preparados para iniciar el proceso de transformación. Se selecciona la tabla de datos a transformar en una tabla con valores de verdad desde una ventana. Usted establece el corte o acepta el predeterminado y luego hace clic en "OK". La siguiente ilustración muestra la tabla de resultados con los valores de verdad "Example_DT_TT.csv" para la tabla de datos "Example_DT.csv" (_TT significa Truth Table/Tabla de Verdad y se añade automáticamente al nombre de la tabla de datos):



10.3 Determinación de los implicados en los criterios "positivos" y "negativos"

Con los valores de verdad ahora podemos determinar finalmente los implicantes. El submenú nos enfrenta a la decisión de si queremos escoger todos los casos en los que la condición a seleccionar como criterio es "positiva", es decir, "verdadera" - o si queremos determinar los implicantes para aquellos casos en los que el criterio es "negativo", es decir, "falso". Como veremos, ambas formas de minimizar las configuraciones condicionales se complementan de forma muy útil. Quedémonos primero con la opción "Criterio: VERDADERO".

Después de haber seleccionado el nombre de la tabla de valores de verdad con la que queremos realizar la minimización, la tabla se carga en la pantalla. Ahora tenemos que seleccionar una de las condiciones (¡columnas!) como criterio haciendo clic en la cabecera de la columna. A continuación, AQUAD busca las combinaciones para las que la condición crítica (nuestro criterio) es "verdadera" a partir de todos los valores de verdad de nuestra tabla.



Le sorprenderá ver que ya no hay 12 casos en la tabla, sino sólo siete combinaciones. AQUAD elimina todas las redundancias de la mesa. Es decir, para los casos con valores de verdad idénticos, sólo queda una combinación en la tabla final de minimización lógica como sustituta de las demás.

En nuestro caso, elegimos la condición "D" como criterio. Es decir, buscamos aquellas configuraciones de condiciones que puedan desempeñar un papel en la aparición de un rendimiento escolar superior a la media (criterio "D") junto con el tamaño de la clase. O dicho de otra manera: En nuestro ejemplo, elegimos la condición "D" como criterio porque queríamos encontrar entre los diversos criterios aquellos que (junto con el tamaño de la clase) desempeñan un papel significativo en los casos en los que observamos el rendimiento de la escuela superior.

El hecho de que el criterio de minimización tenga que determinarse en cada caso señala otra característica especial de AQUAD: al construir la tabla de datos, no hay que decidir desde el principio qué característica va a ser el criterio posterior y, por tanto, colocarla en la última columna, por ejemplo. Esto sólo tendría sentido en relación con los planes de investigación que definen variables independientes y dependientes.

En cambio, para muchas preguntas de la investigación cualitativa, tales presuposiciones sobre los vínculos causales entre las condiciones son la excepción. AQUAD permite elegir cualquier categoría (código, rasgo, significado) del total de categorías utilizadas y buscar las configuraciones típicas de los demás rasgos que se dan junto a la expresión "verdadero" (o "falso") de la categoría seleccionada como criterio. Por tanto, la minimización booleana sirve principalmente para fines heurísticos en AQUAD. Veamos ahora el resultado en nuestro ejemplo:

work}} - Editor

Datei Bearbeiten Format Ansicht Hilfe
Minimización booleana - archivo: qca-22_TT.csv
Criterio: Condición 4 / VERDAD

```
Implicante/s
  Bc
  AC
  AB
CASOS:

3 Implicante/s
--> 1. implicante: Bc - casos:
  2  3 12
--> 2. implicante: AC - casos:
  1  4  9 10
--> 3. implicante: AB - casos:
  3  4 10
```

```
Implicante Bc: 3 casos
Implicante AC: 4 casos
Implicante AB: 3 casos
```

```
Implicantes esenciales
--> 1. implicante: AC - casos:
  1  4  9 10
--> 2. implicante: Bc - casos:
  2  3 12
```

```
Implicante AC: 4 casos
Implicante Bc: 3 casos
```

```
Implicantes esenciales
```

En nuestro ejemplo, la minimización para el criterio "D" produce Bc, AC y AB como principales implicados. Es decir, se observa un alto rendimiento en las clases escolares bajo las siguientes condiciones:

- Con un nivel de habilidad alto (B) y un tiempo de observación relativamente corto (c) (Bc representa la relación lógica B y c); o
- con un bajo número de alumnos (A) y un largo tiempo de observación (C); o
- con un número reducido de alumnos ("A") y un nivel de capacidad elevado ("B").

Si tenemos en cuenta el solapamiento de los mapeos de los casos, podemos resumir que en este estudio se observa un alto rendimiento en todos aquellos casos en los que las clases eran pequeñas y la observación tuvo lugar durante un largo periodo de tiempo (AC), o en aquellos en los que había un alto nivel de habilidad con un periodo de observación corto. En el primer caso, la capacidad no desempeña un papel importante; en el segundo, el tamaño de la clase sí.

En configuraciones condicionales complejas, la búsqueda de implicantes esenciales puede no dar un resultado completo. Esto significa que las configuraciones de condiciones no están completamente cubiertas por las configuraciones reducidas.

No sólo obtenemos información sobre las configuraciones de las condiciones, aquí para el criterio "D", sino también información sobre los grupos o clusters de casos comparables. Como muestra también el ejemplo, los principales implicados son a menudo redundantes, es decir, forman grupos de casos superpuestos. Hay casos que pertenecen a dos configuraciones diferentes, concretamente los casos 3, 4 y 10.

¿En qué consiste la opción "criterio negativo"? Como has adivinado correctamente, la única diferencia es que en este caso AQUAD elige una combinación de condiciones para la minimización para la que es cierto que el criterio elegido es "falso" en cada caso. Como resultado, obtenemos las "condiciones negativas" del criterio, es decir, obtenemos aquellos implicantes que están relacionados con la expresión lógicamente "incorrecta" de la condición del criterio.

En nuestro ejemplo, vemos que el bajo éxito escolar (d) se observa en las clases grandes con aptitud moderada (ab) o en las clases con aptitud moderada y con un periodo experimental largo (bc) o en las clases grandes y con un periodo experimental largo (aC). Si hace esta minimización con su tabla de ejemplo, verá que los implicantes principales no tienen que ser necesariamente redundantes. Obtenemos un resultado con grupos superpuestos en este caso; no hay diferencia entre implicantes principales e implicantes esenciales aquí.

10.4 Para qué más pueden servir los implicantes...

De los ejemplos anteriores se desprende que la lógica de minimización de AQUAD se utiliza para comparar análisis cualitativos, de ahí el nombre de "Análisis comparativo cualitativo" para el procedimiento. En particular, se puede utilizar para comparar contextos de significado o configuraciones de categorías

- si está examinando un gran número de casos individuales o
- si quiere realizar un meta-análisis de estudios cualitativos.

Esto hace que los puntos comunes y las diferencias entre los casos o estudios individuales sean claramente visibles. En los últimos pasos, cuando queremos resumir o agrupar nuestros resultados para diferenciar entre tipos de texto o hablantes, el proceso de minimización lógica parece inevitable.

En el esfuerzo por descubrir relaciones causales más allá de los límites de las condiciones específicas de cada caso, necesitamos identificar una categoría en todos los casos que sea el efecto, por así decirlo, en cuyas posibles causas estemos interesados. Por ejemplo, la investigación de Marcelo (1991) estaba interesada, entre otras cosas, en encontrar las causas de las dificultades disciplinarias de los jóvenes profesores en sus clases. En la formulación lógica "si... entonces..." de la causalidad empírica, la categoría del efecto define la parte "entonces". En concreto, el análisis abordó el problema "Cuando ocurre algo aún desconocido, los profesores principiantes tienen problemas de disciplina". Lo que nos interesa es el contenido de la parte "si", es decir, los grupos o configuraciones de categorías que juntos causan el efecto crítico. Dado que tales enlaces "si-entonces" se conocen como la relación lógica de "implicación", también decimos que las proposiciones dentro de la parte "si" implican la proposición "entonces". Por lo tanto, llamamos a estas proposiciones causales los "implicantes" de la implicación.

Para ilustrar el procedimiento, volvemos a tomar ejemplos del estudio de Marcelo (1991) sobre profesores principiantes. Como se ha mencionado anteriormente, el autor comprobó que muchos de estos jóvenes profesores hablaban con mucha frecuencia de problemas de disciplina en sus clases, pero no todos mencionaban este problema. En la búsqueda de diferencias críticas que puedan servir para explicar por qué algunos de los jóvenes profesores experimentan dificultades de disciplina, el análisis se centró en seis categorías, focos sustantivos de las entrevistas, a saber, las declaraciones de los profesores sobre:

- A sí mismos,
- B relaciones profesor-alumno,
- C Métodos de enseñanza,
- D problemas de disciplina,
- E motivación de los alumnos y
- F clima de clase.

El análisis de las configuraciones para la condición D (problemas de disciplina) como criterio reveló tres grupos de implicados:

$$D = ABC + ACEF + abcef.$$

De esta reducción podemos aprender que tenemos que distinguir tres grupos de principiantes con problemas de disciplina (D). La interpretación de estas agrupaciones parece muy pertinente para la organización de la formación de los profesores:

- Configuración ABC: Un primer grupo puede caracterizarse por la configuración ABC, es decir, estos profesores reflexionan sobre sí mismos, sobre las relaciones profesor-alumno y los métodos de enseñanza, pero no sobre la motivación de los alumnos y el clima del aula.
- Configuración ACEF: Un segundo grupo, que se puede caracterizar por la configuración ACEF, habla de sí mismo, de los métodos de enseñanza, de la motivación de los alumnos y del clima social, pero no parece reflexionar sobre las relaciones entre profesores y alumnos.
- Configuración abcef: El tercer grupo, caracterizado por la configuración abcef menciona con frecuencia los problemas de disciplina en la entrevista, ipero ninguna de las otras categorías centrales!

Esta forma de análisis implícito proporciona grupos de casos. Por razones teóricas y metodológicas, pero también por razones prácticas, es posible que queramos alejar nuestra mirada de los resultados generales de las diferentes configuraciones condicionales de una categoría crítica y centrarnos en los casos individuales. En otras palabras: Deberíamos volver a las entrevistas -y especialmente a las transcripciones de las entrevistas de los profesores que pertenecen a uno de los subgrupos con problemas de disciplina- y centrarnos en codificaciones muy específicas. El estudio de la lista de casos que se obtiene como resultado anima a comparar permanentemente también en este nivel de análisis y abre el camino desde un alto nivel de abstracción hasta las formulaciones concretas de los casos.

Además de ayudar a resumir los resultados, la minimización lógica proporciona importantes funciones heurísticas en las primeras fases del análisis. Al analizar los implicantes de las configuraciones condicionales, podemos obtener valiosas pistas sobre la dirección que puede tomar nuestra interpretación. Supongamos que nuestro estudio incluye 50-60 entrevistas. Además, a lo largo de las 10 primeras entrevistas, desarrollamos cinco categorías importantes para la interpretación. Aquí llamamos simplemente a estas categorías A, B, C, D y E. Estas categorías ponen en marcha conjeturas interesantes, pero desgraciadamente contradictorias, sobre importantes relaciones de significado en nuestros datos. Probablemente no querrá continuar con sus esfuerzos interpretativos, codificando texto tras texto, si empieza a dudar de su enfoque en esta fase. ¿Quién querría acabar, después de mucho trabajo, dándose cuenta quizás de que algo crucial se pasó por alto desde el principio?

En su lugar, puede tomar la codificación de las primeras diez o doce entrevistas e identificar una categoría especialmente importante como criterio "positivo", luego como criterio "negativo" de minimización lógica, y hacer que AQUAD encuentre los implicados de ese criterio. Asumiendo que la condición A es el criterio crítico, primero encontramos los implicados para todos aquellos casos en los que la categoría A fue considerada importante o "verdadera" para los entrevistados. De este modo, aprendemos las configuraciones de condiciones de B, C, D y E que pertenecen juntas (tal vez incluso son la causa) para la ocurrencia de la condición A. Las configuraciones resultantes podrían ser las siguientes:

$$A = BD + BC + bcd$$

A continuación, activamos la opción "Criterio negativo" para encontrar configuraciones de condiciones de B, C, D y E en las entrevistas en las que el criterio A no se mencionó nunca o se consideró poco importante, es decir, "falso". Supongamos que ahora encontramos las siguientes configuraciones:

$$A = BD + cde$$

Es evidente que hay una contradicción. La configuración BD se encuentra tanto como una configuración condicional para los enunciados bajo la condición crítica $A = \text{verdadera}$ como para los enunciados bajo la condición

a = falsa. Al cabo de poco tiempo, es decir, después de haber interpretado 10 de las 60 entrevistas, obtenemos así una importante pista heurística sobre cómo diferenciar nuestro sistema de codificación. Probablemente nos olvidamos de incluir los aspectos evaluativos al utilizar las categorías B y C. Supongamos que nuestras entrevistas versan sobre las experiencias de los estudiantes de magisterio durante las prácticas en el aula. Por ejemplo, codificamos las declaraciones sobre sus observaciones relativas a las interacciones entre profesores y alumnos, pero omitimos registrar en nuestros códigos si un estudiante de magisterio experimentaba una interacción concreta como exitosa/positiva o infructuosa/negativa.

Así, en la mayoría de los textos de las entrevistas se encuentran códigos referidos a la categoría B o D, pero sin ninguna referencia a la expresión cualitativa de A. Sin embargo, si ahora tenemos en cuenta si la observación de una interacción fue evaluada como positiva por los entrevistados y, por lo tanto, probablemente se ajusta a la condición A, o si una interacción fue evaluada negativamente y, por lo tanto, no se ajusta a la condición A en absoluto (pero sí a), entonces podemos resolver la contradicción en poco tiempo. Como se puede ver, como herramienta heurística, el análisis de implicantes puede facilitar la labor de descubrir las categorías adecuadas, incluso cuando sólo se han analizado unos pocos textos.

Por último, debemos pensar también en las posibilidades meta-analíticas que nos abre la minimización lógica. Los meta-análisis no tienen que ser necesariamente cuantitativos, aunque Glass, McGaw y Smith (1981) declaran que esto es "innegable". Dependiendo del tipo de datos, es posible realizar meta-análisis cualitativos o cuantitativos, y los dos enfoques pueden no diferir en el rigor y la sistemática de la comparación de los resultados de los estudios. Con una herramienta como AQUAD, este requisito puede cumplirse para los meta-análisis cualitativos.

Las fuentes de error en la comparación de estudios, que Jackson (1980) enumeró en una aleccionadora recopilación, no pueden eliminarse por completo ni siquiera con un soporte informático para el análisis. Como se ha subrayado anteriormente, el investigador realiza el análisis, el ordenador sólo sirve de herramienta. Si, como señala Jackson (1980), sólo un pequeño porcentaje de los autores que utilizan resúmenes de estudios anteriores no los discuten también de forma crítica, la herramienta no sirve de nada. Por otra parte, con el apoyo de la informática, los propios autores deberían darse cuenta de si sólo han recogido constelaciones específicas de condiciones a partir de las conclusiones de otros estudios o análisis de casos o han pasado por alto configuraciones contradictorias. La acusación de Jackson (1980, p. 459) de que la mayoría de las comparaciones se llevan a cabo "con mucho menos rigor del que es posible en la actualidad" adquiere especial relevancia con la disponibilidad de programas informáticos como AQUAD.

10.5 Función de los implicantes en el proceso de construcción teórica

¿Cómo deben concebir los investigadores a los implicados como resultados de los análisis configurativos? ¿Sirven los implicados como pruebas que ayudan a probar o rechazar las teorías que los sustentan, como planos para reconstruir las visiones del mundo de otras personas o para formular teorías que afloran en el proceso de análisis cualitativo? La respuesta, una vez más, es "ni lo uno ni lo otro...". ni". Dado que no es éste el lugar para elaborar consideraciones metodológicas, las explicaciones de Ragin (1987) sobre el diálogo entre las pruebas y las ideas en el análisis configuracional booleano podrían ser de interés como lectura adicional:

El camino booleano hacia la comparación cualitativa ... es un camino intermedio entre dos extremos, el orientado a las variables y el orientado a los casos, es el camino intermedio entre la generalización y la complejidad. Permite a los investigadores hacer ambas cosas, incluir muchos casos y estimar la complejidad causal (Ragin, 1987, p. 168).

Capítulo 11: Cómo combinar los análisis cualitativos y cuantitativos

11.1 Sobre el debate acerca de los métodos de investigación cualitativos frente a los cuantitativos

En la década de 1980, la revista *Educational Researcher*, de la Asociación Americana de Investigación Educativa, dedicada a cuestiones de carácter fundamental, presentaba como tema recurrente los argumentos sobre "calidad" y "cantidad". Las contribuciones, en gran medida interrelacionadas, se leen a menudo como una novela por entregas, aunque los argumentos se agoten pronto. Se pueden distinguir cuatro tipos de textos:

- (1) Textos que intentaron fundamentar una oposición irreconciliable de calidad y cantidad en la investigación educativa (Smith, 1983; Smith & Heshusius, 1986);
- (2) textos que articulan la consternación ante la polarización e iluminan los malentendidos subyacentes (Shulman, 1981; Phillips, 1983; Tuthill & Ashton, 1983; Howe, 1985, 1988; Firestone, 1987) o desarrollan interpretaciones alternativas (Eisner, 1981, 1983);
- (3) textos que respondían a la tesis de la incompatibilidad con estrategias de investigación que incorporaban la calidad y la cantidad de forma complementaria (Miles & Huberman, 1984, 2ª ed. 1994; Huberman, 1987), y
- (4) textos que se centraron en los enfoques cualitativos y trataron de contribuir a su sistematización (Fetterman, 1982, 1988; Jacob, 1988).

Sin embargo, un debate sobre los objetivos y las condiciones de la investigación cualitativa, especialmente en comparación con los enfoques cuantitativos, no parecía querer despegar realmente, al menos en el ámbito de la investigación de orientación psicológica. Uno tiene la impresión de que no se están elaborando tanto aplicaciones complementarias en una reflexión objetiva sobre las ventajas y los problemas de los enfoques cualitativos y cuantitativos, sino que está en marcha un proceso de polarización. Incluso en los grupos de investigación temáticamente especializados, se puede observar la formación de "campos metodológicos" opuestos. En este proceso, la delimitación en cuanto al tipo de datos reconocidos también pareció afectar a los objetivos de la investigación. En ocasiones, esto llevó a que las preguntas y los resultados de los "otros" ya no se tomaran en cuenta y, desde luego, no se discutieran; uno se conformaba con el intercambio de etiquetas - por ejemplo, investigación "blanda" aquí, investigación "dura" allá. Según las posiciones de antes los enfoques cualitativos y cuantitativos se concentran en aspectos opuestos de los fenómenos sociales:

<i>Métodos cualitativos</i>	<i>Métodos cuantitativos</i>
comprensión descripción interpretativo subjetivo emic (perspectiva interna del sujeto)	comprobación pronóstico empírico objetivo etic (perspectiva del observador/científico)

No cabe duda de que las diferentes orientaciones metodológicas no definen contradicciones incompatibles. Hay suficientes ejemplos convincentes de la práctica de la investigación en los que se ha podido reconstruir el mundo subjetivo de las personas investigadas y reducir al mismo tiempo esta información de forma objetiva. Cada enfoque de investigación construye una visión del mundo específica, sólo que algunos métodos dan la impresión de que con ellos se puede encontrar simplemente la realidad. Sin embargo, la estrategia de delimitación no parece haber sido superada ni siquiera en la actualidad. Claramente, esta posición fue expresada en su momento por Smith & Heshusius, 1986, p. 11):

"... si un investigador cuantitativo no está de acuerdo con un investigador cualitativo, después de todo es posible que hable uno con otro? La respuesta está, de todos modos en la actualidad, un no categórico! Una exhortación que hay que aceptar un resultado particular porque está basado sobre hechos no va a influir en una persona que cree que no pueden existir hechos sin interpretación de este caso. Por otro lado la idea que hechos están cargados de evaluaciones y la idea que no hay tribunales de apelación más allá que diálogo y persuasión por lo menos parecerá no científica e insuficiente a un investigador cuantitativo."

La dicotomía y la devaluación mutua implícita parecen insostenibles. Sean cuales sean los métodos utilizados, la investigación debe llevarse a cabo de forma sistemática y conducir a resultados fiables y válidos. E independientemente de que se cuenten las frecuencias de comportamiento o se interpreten las respuestas verbales, deben hacerse visibles las estructuras de significado que el investigador construye para garantizar los criterios mencionados. Ahora bien, no cabe duda de que existen reglas precisas para ello en el ámbito de la metodología cuantitativa, mientras que para algunos investigadores de orientación cualitativa parece tratarse de una cuestión de si hay que revelar el sistema de los propios procesos de construcción de tal manera que uno mismo u otros puedan controlar y, en su caso, repetir estos procesos. En nuestra opinión, este es un requisito indispensable.

Desgraciadamente, los libros de texto de metodología cualitativa, aunque describen detalladamente los más diversos métodos de recogida de datos, tratan los problemas y los procedimientos de análisis de datos de forma bastante madura. Sieber (después de Miles, 1983) observó hace más de 20 años que en los principales libros de texto sólo se dedicaba entre el 5 y el 10% del ámbito al análisis de los datos; hoy la proporción parece ser más favorable, pero sigue siendo muy desequilibrada. Es aquí donde el desarrollo de software de datos cualitativos puede suponer un gran avance. Los ordenadores, como instrumentos de análisis de datos, pueden apoyar los procesos de reducción y análisis de datos cualitativos de manera que se pueda controlar, reconstruir y, sobre todo, comunicar con precisión el curso y los resultados de la investigación. Esto podría iniciar un diálogo entre los miembros de los dos campos, en la línea de cómo Miles & Huberman (1984, p. 23) describieron su propia posición. Para ellos, es importante que los investigadores procedan sistemáticamente y lleguen a un consenso sobre el contenido y la metodología. Ellos mismos se sentían investigadores cualitativos "de derechas" o "positivistas afeminados", ya que recomendaban tanto prestar atención a la realidad en los enfoques interpretativos como, a la inversa, tener claro qué es lo que uno considera realmente "datos" y qué papel juegan las propias perspectivas en el proceso de atribución de significado a los datos.

A este respecto, cabe mencionar de nuevo el papel del análisis de datos asistido por ordenador, formulado de forma algo diferente: Los ordenadores pueden ser útiles en el proceso de investigación cualitativa, ya que ayudan a la comunicación, la reconstrucción y el control de los procesos de investigación. De este modo, los ordenadores pueden estructurar los enfoques naturalistas y ayudar a sistematizar las interpretaciones rigurosas. Al mismo tiempo, el uso de ordenadores apoya el objetivo, muy pragmático pero no menos importante, de simplificar los elaborados procesos de trabajo.

Eisner (1981, p. 6) consideraba que esa concepción de los métodos limitaba demasiado las posibilidades y las tareas de los métodos cualitativos. Quería dar más peso a los procesos "artísticos" e intuitivos: "El comportamiento

manifiesto se trata principalmente como una pista, como un peldaño para llegar a otro lugar. La otra posibilidad es la "habitabilidad", la empatía; es decir, participar imaginativamente en la experiencia de los demás. ... La diferencia es sutil pero significativa. En la primera elección de objeto, el observable se trata de forma estadística; se estima intuitivamente (o estadísticamente) la probabilidad de que este comportamiento tenga uno u otro significado particular. ... Esta última opción se basa en la capacidad ... de proyectarse intuitivamente en la vida de otro para experimentar lo que esa persona está viviendo. ... Un punto central, por tanto, en los enfoques artísticos de la investigación son los significados y las experiencias de las personas que operan en la red cultural objeto de estudio".

Afirmaciones como las de Eisner (1981) se malinterpretan fácilmente (y de buen grado) si se supone que pretenden, en un dogmatismo inverso, sustituir los hechos por la ficción. Por el contrario, la complementariedad mutua de las perspectivas de la investigación cualitativa y cuantitativa, determinada también por la elección del tema, queda ilustrada muy vivamente. Por poner un ejemplo, que para los lectores que juran la investigación cuantitativa puede que ya exceda los límites de lo científicamente admisible: ¿Qué conocimientos transmiten los estudios cuantitativos sobre la situación de los niños discapacitados mentales y qué se aprende en comparación, por ejemplo, con la lectura de "Hirbel" de Härtling? Pero ¿de qué sirve empatizar con este destino si no se tiene una base segura para sacar conclusiones en otras situaciones? La cantidad por sí misma no tiene sentido, la calidad por sí misma sigue siendo intrascendente. Especialmente la investigación en campos orientados a la aplicación, como la ciencia de la educación o la psicología de la educación, depende de la complementación inteligente de los enfoques cualitativos y cuantitativos de sus áreas temáticas.

Por lo tanto, uno no puede sino sacudir la cabeza con desconcierto cuando lee la declaración de Smith y Heshusius (1986, p. 11) citada anteriormente como "conclusión del debate". Sin embargo, no hay que limitarse a sacudir la cabeza. Las consecuencias de esta dicotomización son demasiado graves para que las ciencias sociales no intenten salvar el abismo que se ha abierto aquí. Una forma de hacerlo parece ser exigir que, cuando se utilicen métodos cuantitativos, se abran y revelen cualitativamente las dimensiones de significado de los datos, y cuando se utilicen métodos cualitativos, se sistematicen y documenten los procesos de interpretación y, además, se ordenen los hallazgos de forma cuantitativa. Si se cumple este requisito, se puede y se debe recurrir a los resultados y procedimientos de la otra orientación metodológica en el marco de cada una de las dos orientaciones aparentemente incompatibles, y hacerlo en beneficio de la propia investigación. El requisito previo para ello es la transparencia y la regularidad, la verdadera metodología en la aplicación práctica de los enfoques metodológicos en la investigación.

En este sentido, no existe una investigación "buena" o "mala", "dura" o "blanda", sino sólo rigor metodológico o dejadez – en ambos enfoques –, independientemente de la calidad del anclaje teórico de un estudio. Para la apertura y aceptación de los enfoques cualitativos, el desarrollo de procedimientos sistemáticos para el análisis de datos me parece prometedor.

En el debate sobre las posibilidades del análisis de contenido cualitativo y en el desarrollo de procedimientos practicables, existen afortunadamente muchas propuestas sobre cómo se puede tener en cuenta sistemáticamente la relación complementaria entre calidad y cantidad. Ya en 1988, Mayring propuso un modelo global de fases para optimizar la relación entre el análisis cualitativo y el cuantitativo. Según este modelo, el análisis de contenido parte de los determinantes cualitativos de la pregunta de investigación, los términos o categorías críticas y los instrumentos. Según el objetivo y/o el tema del análisis, los instrumentos de análisis pueden utilizarse recurriendo a procedimientos cuantitativos (por ejemplo, determinación de las frecuencias de las palabras, análisis de conglomerados de términos críticos). En la fase final, los resultados de este paso se relacionan con la pregunta de investigación y se extraen conclusiones; esto también requiere un procedimiento cualitativo-interpretativo. Villar y Marcelo (1992) demuestran y discuten diferentes formas de combinar los métodos cualitativos y cuantitativos en el diseño de los estudios, en la recogida de datos y en el análisis de los mismos, utilizando su propia investigación como ejemplo.

Este esquema nos parece que crea el excelente prerrequisito para cumplir la exigencia de Lisch y Kriz (1978, p. 46), que también se hizo hace tiempo, de no eliminar las partes subjetivas en el análisis de contenido, sino de hacerlas explícitas: "Haciendo que el analista de contenido sea lo más consciente posible de sus decisiones en la reconstrucción de la realidad social y comunicándolas a su comunidad científica el marco interpretativo se vuelve intersubjetivamente comprensible y verificable y, por lo tanto, se puede decidir la cuestión de la clasificación de los resultados analíticos del contenido en una teoría relevante para la acción. ... La experiencia específica del analista de contenidos con un texto se hace comunicable, reconstruible y, por tanto, en la medida de lo posible, reexperimentable".

Tanto los desarrollos metodológicos de los últimos años como la práctica de la investigación en las ciencias sociales – al menos fuera de Alemania – muestran ahora, afortunadamente, que la investigación se orienta cada vez más hacia los retos de las respectivas preguntas de investigación en lugar de hacia el marco de los métodos específicos. El uso opcional de métodos cualitativos y cuantitativos según una estrategia de "metodología mixta" (Mathison, 1988; Tashakkori y Teddlie, 1998) promete permitir trabajar sobre aquellos aspectos de un problema, de una pregunta de investigación compleja, que no podrían haber sido investigados sin esta orientación debido al repertorio de métodos disponibles en cada caso - o que no habrían sido percibidos en absoluto debido a las anteojerías metodológicas. El hecho de que también se haya publicado un manual sobre enfoques y aplicaciones de los "métodos mixtos" (Tashakkori & Teddlie, 2003) sugiere un desarrollo fructífero de este enfoque en la práctica de la investigación en ciencias sociales y del comportamiento. Para revisiones recientes, véase Mayring, Huber, Gürtler y Kiegelmann (2007) y Gläser-Zikuda, Seidel, Rohlf, Gröschner y Ziegelbauer (2012).

11.2 "Métodos mixtos": ¿Palabra de moda o estrategia de investigación?

Así pues, los métodos cualitativos y cuantitativos no se excluyen mutuamente, sino que tienen una relación complementaria. En primer lugar, es cierto que la elección del método depende de la pregunta a la que se pretende responder con un estudio. Si quiere estimar cuántas personas votarán a un determinado partido en unas elecciones o quieren comprar un determinado producto en el transcurso de los próximos seis meses, le interesan los valores numéricos más fiables posibles como resultado. Si, por el contrario, quiere averiguar por qué la gente quiere comprar un determinado producto o por qué prefiere un producto de la competencia, preferirá las evaluaciones diferenciadas desde el punto de vista subjetivo de los clientes.

Si se observan detenidamente los dos estudios reseñados, en los que en el primer caso hay un valor porcentual con indicación de los límites de confianza, es decir, una indicación de cantidad, y en el segundo caso una recopilación de cualidades atribuidas, se observa que en ambos estudios se producen actividades cualitativas y cuantitativas. Sin embargo, en el primer caso no se suele discutir el papel de los procedimientos cualitativos, y en el segundo el de los cuantitativos.

El punto de partida de la investigación es un problema, en el ámbito de las ciencias sociales normalmente un problema en un ámbito de la vida o campo de acción social. La percepción más o menos difusa de este problema toma forma en la formulación de una o varias preguntas de investigación. En estas fases no se comprueban las hipótesis de forma cuantitativa, sino que primero se especifican las preguntas. En la exploración posterior, ahora hay que formular hipótesis en los estudios cuantitativos y redactar un diseño de investigación, así como seleccionar métodos con los que sea posible la comprobación (cuantitativa) de las hipótesis. En el caso de los estudios cualitativos, el diseño se centra en las posibilidades metodológicas para obtener más información sobre el área del problema de forma que se puedan formular hipótesis científicas sobre ella. Tanto si se trabaja con tests, escalas de valoración, cuestionarios, entrevistas, revisión de diarios, etc., como si se trata de un método de recogida de datos, en cualquier caso es necesario tener

acceso al campo en el que se puede recoger la información. Establecer contactos sobre el terreno, generar confianza y crear la voluntad de participar en el estudio son actividades de investigación necesarias más allá de la visión polarizadora de los métodos cualitativos y cuantitativos.

Sólo queremos mencionar, pero no discutir más, el hecho de que, por supuesto, los instrumentos cuantitativos utilizados en los estudios cuantitativos de la fase siguiente para explicar el problema no fueron desarrollados en sí mismos exclusivamente con métodos cuantitativos. Incluso en la construcción de un test altamente estandarizado, alguien, en algún momento, debe haber tenido la idea y decidido que los ítems del test pueden captar lo que el test debe medir. En general, es importante recordar que los propios investigadores y los papeles que se atribuyen en sus empresas a partir de sus orientaciones epistemológicas determinan la gama de actividades de investigación que se reflejan como significativas en un contexto de investigación determinado. Maxwell (2002) destaca la influencia de las características personales del investigador, por un lado, por ejemplo, las experiencias previas, las convicciones, los objetivos personales, y por otro lado, las relaciones personales del investigador con su campo de investigación, es decir, principalmente con el investigado. La forma en que los investigadores conceptualizan y se comprometen con un estudio depende de su identidad y perspectiva, así como de sus relaciones con los sujetos de la investigación. A la inversa, esto también se aplica a los investigados.

Sea cual sea el enfoque metodológico elegido, ahora hay que recoger los datos empíricos con los instrumentos elegidos y analizarlos. Hay que comprobar la validez de los resultados; también en este caso, según el criterio de validez, los enfoques cualitativos y cuantitativos se complementan. Por último, hay que redactar un informe de investigación. Los lectores del informe y/o los propios investigadores considerarán entonces si los resultados pueden contribuir a la solución del problema investigado y de qué manera, y al hacerlo tampoco se guiarán por la dicotomía de los métodos cualitativos frente a los cuantitativos.

Incluso si se está de acuerdo en principio con estos análisis y se admite la complementariedad mutua de los enfoques cualitativo y cuantitativo en el nivel abstracto del contexto general de las fases del proceso de investigación, la cuestión sigue siendo si la metodología cualitativa y cuantitativa puede combinarse sistemáticamente en el nivel concreto de las secciones individuales, especialmente la recogida y el análisis de datos. Más adelante veremos que tales combinaciones ya están implícitas en el proceso de investigación cualitativa, previstas explícitamente como una característica importante del diseño en la planificación, realización y validación de estudios completos.

11.2.1 Combinación implícita de métodos en los procesos de investigación cualitativa

Anteriormente describimos la interacción de las estrategias de diferenciación y generalización de la codificación con referencia a Shelly y Sibert (1992) como un proceso cíclico de inferencia inductiva y deductiva. Mayring (2001) señala la generación sistemática de categorías y la asignación de segmentos de datos como un objetivo común de ambos procesos y concluye: "Cuando se trabaja sistemáticamente con categorías de esta manera, tiene sentido concebir estas asignaciones como "datos" y trabajar con ellas cuantitativamente en un segundo paso del análisis" (párrafo [16]).

Así, los resultados de una primera reducción cualitativa-interpretativa de los datos iniciales se cuentan y ordenan por frecuencia, se expresan en porcentajes, se ordenan en escalas ordinales (mucho - medio - poco), se comparan entre segmentos de expedientes y/o casos con procedimientos estadísticos, etc. La contribución de estos nuevos resultados para responder a las preguntas de la investigación debe interpretarse de nuevo cualitativamente en un tercer paso (Mayring, 2001, [17]).

11.2.2 Combinación explícita de métodos en los procesos de investigación cualitativa/cuantitativa

Según Villar y Marcelo (1992), el progreso de la investigación en ciencias sociales se basa precisamente en que se hace uso de una variedad de métodos según el problema de investigación - y no se limita el problema y los enfoques empíricos según las posibilidades de un enfoque metodológico. Se refieren a la sugerencia de Greene, Caracelli y Graham (1989) de recopilar datos en un diseño de "método mixto". Como señalan Villar y Marcelo (1992), esto abre el proceso de investigación a la exigencia de "triangulación", que al menos aumenta la validez de la investigación al incluir diferentes métodos, fuentes de datos e investigadores (Mathison, 1988).

En Villar y Marcelo (1992), esta idea se desarrolla concretamente para las fases centrales de la investigación empírica, a saber, el acceso al campo, los métodos de recogida de datos y el análisis de los mismos. Los autores muestran procedimientos concretos que ellos mismos han utilizado en sus estudios para que las ventajas de los métodos cualitativos y cuantitativos fructifiquen en su trabajo de forma complementaria:

- Combinación de métodos en la fase de acceso al campo.

Cuando en un estudio es importante trabajar con una muestra representativa de una población definida (por ejemplo, "los" profesores de primaria; "los" profesores de matemáticas de secundaria superior; etc.), se suele combinar una selección de sujetos basada en la cantidad con un control de la representatividad basado en las características cualitativas de la muestra así extraída. Como ejemplo, Villar y Marcelo (1992) citan un estudio sobre la socialización de profesores noveles en el que los participantes fueron seleccionados de una lista mediante una tabla de números aleatorios. A continuación, examinaron la composición de la muestra en función del sexo de los sujetos, su asignatura, el tipo de centro y la ubicación sociogeográfica del mismo, todas ellas características cualitativas.

Pero también en el caso de los estudios de casos individuales, más habituales en la investigación cualitativa, la "mezcla de métodos" en esta fase puede aumentar la precisión de la selección de los casos individuales y contribuir así, por un lado, a la utilización óptima de los escasos recursos de un proyecto de investigación y, por otro, a aumentar la validez de los resultados. En el "muestreo teórico", los casos no se seleccionan de forma aleatoria, sino específicamente de manera que, según el diseño del estudio, sean típicos o críticos para la(s) pregunta(s) de la investigación, extremos o únicos, o simplemente particularmente informativos porque pueden no haber sido accesibles antes (Yin, 1989). Las herramientas de diagnóstico cuantitativo pueden ser ventajosas en este caso: Si, para un estudio cualitativo sobre los efectos subjetivos y los procesos de aprendizaje estimulados individualmente al utilizar determinados métodos de enseñanza, sospechamos que los alumnos de alto rendimiento y los más débiles se expresarán de forma diferente, podemos utilizar los resultados de las pruebas de rendimiento escolar y/o de inteligencia adecuadas para especificar la selección de sólo unos pocos sujetos a los que podamos entrevistar u observar por razones prácticas.

Estos ejemplos muestran, como también debería expresarse en el modelo de proceso del diagrama anterior, que las decisiones se superponen y entrelazan las distintas etapas del proceso. Así, por supuesto, las decisiones sobre los métodos e instrumentos empíricos se toman en la fase de diseño, que luego tienen un impacto práctico en la fase de recogida de datos.

- Combinación de métodos en la fase de recogida de datos

También aquí hay que partir de la fórmula: "Depende...". - en la pregunta de investigación, es decir, si se quiere juzgar qué combinación se debe recomendar. Para el ámbito de la investigación de orientación cualitativa sobre los procesos de enseñanza/aprendizaje en la formación inicial y continua del profesorado, Villar y Marcelo (1992) citan la observación participante, la entrevista en profundidad y el diario de clase como métodos de recogida de datos frecuentemente utilizados. En cambio, los enfoques de la investigación de impacto en este ámbito suelen recurrir a la observación sin participación y determinan la frecuencia de ciertos fenómenos según sistemas de categorías predefinidos.

Como ejemplo de una de las combinaciones posibles, los autores describen un estudio cuantitativo en el que se registró la práctica docente en la Universidad de Sevilla (Villar, 1981) mediante observación no participante (evaluación de grabaciones con sistemas de análisis predefinidos), se entrevistó a los participantes para obtener información sobre sus puntos de vista subjetivos y luego se les pidió que completaran escalas de valoración (cuantitativas) sobre sus creencias, problemas y el clima social del seminario. Con esta combinación, la complementación mutua de los diferentes hallazgos en el sentido de una profundización de la información de la observación a través de la información subjetiva en la entrevista es tan posible como una triangulación (metódica) mediante la comparación de los datos de las entrevistas y las escalas de valoración (véase más adelante, así como Mayring, 2001).

- Combinación de métodos en la fase de análisis de datos

Las combinaciones de técnicas de análisis cualitativo y cuantitativo son posibles tanto en la evaluación de los datos procedentes de un enfoque metodológico específico, por ejemplo, los datos de las entrevistas, como en la evaluación de los datos obtenidos con métodos diferentes, por ejemplo, en la combinación de entrevista y cuestionario.

Un ejemplo de la primera situación ya se ha descrito en detalle anteriormente. Un simple análisis de frecuencia ayudó a identificar los contenidos críticos de un gran número de entrevistas. Posteriormente, se siguieron utilizando los análisis de correlación para comprobar las correlaciones entre los enunciados críticos (cf. Villar y Marcelo, 1992).

El segundo tipo de situación es característico de los análisis de triangulación. Por un lado, las aportaciones específicas de los distintos métodos pueden utilizarse en el análisis desde una perspectiva cualitativa (cf. Medina, Feliz, Domínguez y Pérez, 2002; véase más adelante); por otro lado, los análisis cualitativos y cuantitativos, especialmente los procedimientos de análisis de dimensiones, como los análisis factoriales o de conglomerados y los procedimientos de regresión, pueden prestar buenos servicios para identificar o confirmar patrones centrales de argumentación. Haciendo hincapié en el enfoque cuantitativo, Schweizer (2004) describe detalladamente, utilizando el primer capítulo del "Quijote" de Cervantes como ejemplo, cómo los hallazgos hermenéuticos relevantes (por ejemplo, aquí sobre la importancia de la lectura para el desarrollo individual, así como sobre las posibilidades de acceso a los motivos e intereses del autor real mirando detrás de su mundo ficticio) pueden elaborarse a través del análisis (cuantitativo) de las formas de palabras/vocabulario y la evaluación de su contribución a la comprensión del texto.

En una última fase del análisis de los datos, a la que nos referimos en el contexto de Aquad con Ragin (1987) como "análisis implícito" o "minimización lógica" o "análisis comparativo cualitativo", el objetivo es generalizar las conclusiones interpretativas más allá del conjunto de datos individuales o del caso. Si preguntamos qué es lo general o típico de un caso concreto, obtendremos como respuesta detalles de las características relevantes que tiene en común

con una clase de casos similares. Por el contrario, si nos preguntamos qué es lo que hace que un caso parezca único, la respuesta nos indicará los atributos que lo distinguen de los demás. Sólo nos satisface la clasificación de los casos en clases cuando cada una de ellas contiene sólo casos que son lo más parecidos posible entre sí en cuanto a atributos críticos, pero que al mismo tiempo difieren lo más posible de los casos de otras clases. Con el proceso de clasificación, el objetivo no es sólo un orden y una representación de los casos actuales, sino que en este orden se intenta aclarar implícita o explícitamente las conexiones, establecer vínculos entre las experiencias, construir expectativas, hacer posibles las predicciones. Se pueden encontrar ejemplos de ello en el capítulo 11 y en Huber (2001). Mayring (2001) ve en la "designación del caso individual como típico para una determinada área temática ... un primer paso de generalización cuantificadora..." [18].

11.2.3 Combinación de métodos en el diseño de la investigación

A nivel de diseño de estudios empíricos, Mayring (2001) ha propuesto cuatro modelos de combinación de métodos cualitativos y cuantitativos, que denomina modelo de preestudio, modelo de generalización, modelo de profundización y modelo de triangulación. Asignamos estas combinaciones a dos tipos y las complementamos con un tercer tipo de combinación:

(1) Macrosecuencias de combinación de métodos:

- El modelo de pre-estudio:

Este modelo describe el enfoque según el cual se postulaba una especie de división del trabajo en la "comprensión científica empírica tradicional", "... según la cual los métodos cualitativos se utilizan principalmente al principio de un proyecto de investigación en el sentido de estudios exploratorios ... se utilizan. Tienen por objeto aclarar el campo de la investigación, diferenciar los problemas y las preguntas, establecer hipótesis preliminares y comprobar la viabilidad de los métodos de investigación" (Krapp, Hofer y Prell, 1982, p. 130). Esta visión de la mera asistencia a la investigación cuantitativa "real" no hace justicia a la función complementaria de los enfoques cualitativos en la investigación empírica. Sin embargo, este modelo deja claro que los estudios cuantitativos difícilmente pueden prescindir de una base cualitativa. Esquemáticamente, la conexión en el modelo de estudio preliminar tiene el siguiente aspecto (cf. Mayring, 2001 [21]):

- El modelo de generalización:

Según este enfoque combinado, se lleva a cabo un auténtico estudio cualitativo independiente, cuyos resultados empíricos ya se atribuyen importancia por derecho propio. Sin embargo, no se quiere o no se puede -sobre todo en el contexto de la aplicación (fase de aplicación; véase más arriba) – sacar conclusiones significativas para las medidas de cambio de la práctica a partir de la base de datos relativamente estrecha de los estudios cualitativos hasta que se haya comprobado cuantitativamente la generalizabilidad de los resultados y se haya asegurado estadísticamente su probabilidad o la fuerza de los efectos de las nuevas medidas. Mayring (2001) da el ejemplo de un análisis de casos cuyos resultados se verifican en un estudio de seguimiento representativo:

- El modelo en profundización:

Las etapas del diseño se invierten en este modelo, es decir, a una encuesta cuantitativa representativa y al análisis de datos le sigue un estudio cualitativo más pequeño, centrado en el caso, cuyos resultados deberían ayudar a comprender mejor la importancia de los resultados cuantitativos (cf. Huber, A. A., 2007).

Muchos estudios del grupo de investigación sobre la juventud de J. Held en el Departamento de Psicología de la Educación de la Universidad de Tubinga están diseñados según este modelo (cf. Held, 1994; Kiegelmann, Huber, Held y Ertel, 2000). En un proyecto sobre "Orientaciones políticas de los jóvenes trabajadores", Held, Horn y Marvakis (1996) investigaron las razones de los jóvenes para sus orientaciones. En el primer paso, los autores presentaron un cuestionario a los jóvenes, en el segundo paso los jóvenes ayudaron a interpretar sus resultados en entrevistas de seguimiento y así profundizar y concretar los resultados cuantitativos.

En este estudio, al igual que en uno anterior de Held y Marvakis (1992), se hace visible otro efecto del modelo de profundización, su vinculación de la explicación y la aplicación, del conocimiento y de la aplicación en el proceso de investigación: La característica del procedimiento metodológico fue que el contacto con los jóvenes entrevistados no terminó con la aplicación de una encuesta por cuestionario. A continuación, se realizó la retroalimentación y el debate de los resultados en las escuelas y las empresas. Los resultados sirvieron de estímulo para las discusiones de grupo y las entrevistas en profundidad, el instrumento cuantitativo se convirtió en un medio de cambio potencial. El proceso de investigación se convierte así en un proceso educativo práctico al mismo tiempo.

A. A. Huber (2007) ofrece un ejemplo detallado de cómo en un estudio sobre el aprendizaje cooperativo se pueden interpretar mejor los amplios resultados cuantitativos y explicar en parte los resultados inesperados.

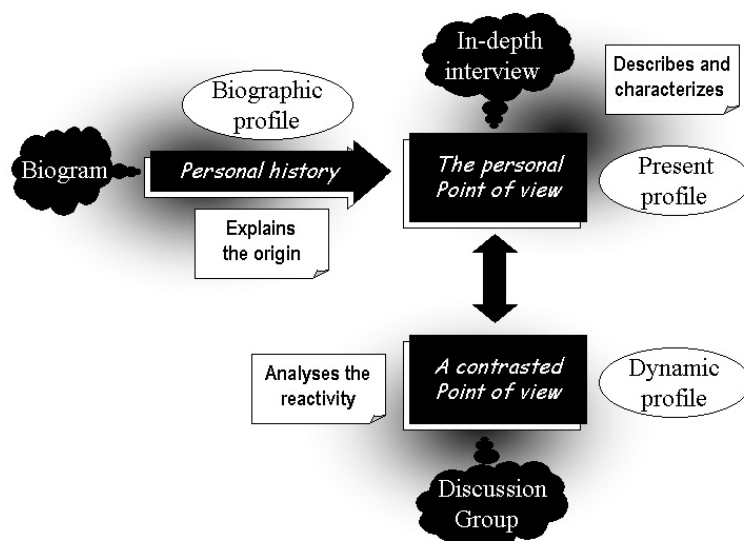
(2) Aplicación simultánea de métodos cualitativos y cuantitativos.

- El modelo de triangulación:

En este enfoque, coexisten simultáneamente y en igualdad de condiciones diferentes enfoques metodológicos en el campo de la investigación. Es importante comparar las respuestas más o menos variadas que resultan de la perspectiva de los diferentes métodos a la pregunta de investigación y determinar la intersección de los resultados individuales como conclusión final.

En el nivel de la triangulación de métodos (véase más arriba), pueden darse todas las combinaciones imaginables, es decir, en el ámbito de la investigación cualitativa, la combinación de métodos cualitativos y cuantitativos o la combinación de diferentes métodos cualitativos. No se utilizan en la secuencia requerida por el diseño, sino simultáneamente o en paralelo. Un ejemplo bien justificado de esta última triangulación lo ofrecen Medina, Feliz, Domínguez y Pérez, (2002, p. 178, original en inglés):

El biograma ayuda a abrir la historia personal o profesional, la entrevista en profundidad contribuye a una comprensión diferenciada de los aspectos individuales relevantes y, por último, la discusión en grupo permite relacionar y comparar las perspectivas subjetivas de los participantes individuales en el estudio. En la intersección de la perspectiva temporal-biográfica, la perspectiva personal-individual y la perspectiva dinámica-social-interactiva,



encontramos finalmente respuestas equilibradas que pueden ser valoradas en su alcance a través de las diferentes perspectivas.

Varias contribuciones en Tashakkori y Teddlie (2003) informan sobre otras variantes del modelo de triangulación, por ejemplo, la combinación de las perspectivas de varios investigadores o diferentes orientaciones teóricas, especialmente el artículo sobre las reglas de integración de Erzberger y Kelle (2003). Sobre la relación entre los "métodos mixtos" y la triangulación, Gürtler y Huber (2012) ofrecen una visión general.

(3) Microsecuencias de combinación de métodos:

Tanto en la combinación de métodos implícitos como explícitos, en función de la decisión de una estrategia concreta, a menudo se plantea la cuestión de cómo se pueden analizar los datos cualitativos disponibles también con la ayuda de métodos cuantitativos. Las respuestas a esta pregunta pueden resumirse en el tercer tipo de combinación, la generación de microsecuencias de combinación de métodos. La forma de conseguirlo es, en cualquier caso, mediante la reducción y la asignación de aspectos seleccionados del material de datos a una escala cuantitativa, normalmente a una escala nominal. Con respecto al procedimiento concreto, esto significa que se cuentan los elementos seleccionados del material de datos, se compilan las frecuencias en listas y, finalmente, estas listas se convierten en forma de tabla para las evaluaciones estadísticas. Los resultados del análisis cuantitativo suelen ser el punto de partida para posteriores interpretaciones cualitativas, por ejemplo, en forma de metacodificación y posterior análisis de implicados.

Ya hemos hablado anteriormente de la determinación de las frecuencias de las palabras; en la siguiente sección se ofrece más información al respecto. En la sección 11.4 también describimos la determinación de las frecuencias de los códigos. Por último, en la sección 11.5 se analizan las microsecuencias de la combinación de métodos con dos ejemplos concretos.

Vemos que las ventajas de combinar datos/métodos cualitativos y cuantitativos pueden utilizarse en el nivel primario de los datos dados, aquí las palabras individuales, pero también en el nivel secundario de los resultados de nuestra interpretación, aquí las codificaciones. A continuación veremos el procedimiento en detalle.

11.3 Frecuencias de palabras

Berelson (1952) desarrolló los fundamentos del análisis de contenido, que para él era un procedimiento cuantitativo basado en la determinación de la frecuencia de los elementos del contenido manifiesto. Ya hemos mencionado que los procedimientos no cuantitativos son ciertamente significativos para determinar y seleccionar los elementos relevantes para el análisis. Concentrémonos aquí en la "mecánica" del proceso.

El acceso a las funciones de recuento en textos se encuentra en el menú principal de Aquad en el grupo de funciones "Análisis de contenido" -> "Análisis de contenido cuantitativo". La siguiente figura ofrece una visión general de cómo se pueden crear listas de palabras críticas. El "recuento de palabras" puede seleccionarse en el submenú anterior de las subfunciones. Los detalles se encuentran en el capítulo 8.

11.4 Frecuencias del código

La combinación de análisis cualitativos y cuantitativos es especialmente interesante con resultados en el nivel secundario de los análisis. El enfoque consiste en contar la frecuencia de los códigos. Se pueden contar todos los tipos de códigos. Especialmente el recuento de los códigos de secuencia, que pueden insertarse adicionalmente en los archivos de códigos con las diversas formas de análisis de "enlaces" como resultado, podría ser interesante en fases avanzadas del análisis.

Para el recuento, el programa necesita una lista de códigos para buscar en los archivos de códigos. Esta lista puede elaborarse de nuevo y actualizarse para cada recuento seleccionando en el registro de códigos. La alternativa sería generar una lista de códigos relevantes una vez por selección y guardarla. Para el recuento, basta con cargar esta lista. A continuación, examinamos con más detalle esta segunda posibilidad. La secuencia de pasos comienza, pues, con la creación de una lista de códigos que consideramos importantes para una pregunta concreta:

① Menú principal: "Búsqueda" -> "Catálogo de códigos" -> "Crear catálogo de códigos".

A continuación, iniciamos la función de recuento de códigos, cargamos la lista que acabamos de crear, iniciamos el proceso de recuento y, finalmente, guardamos los resultados en forma de lista de frecuencias de los códigos seleccionados:

② Menú principal: "Búsqueda" -> "Contar códigos".

Para la exportación y posterior uso de estas frecuencias simplemente enumeradas en programas estadísticos, necesitamos una representación en forma de tabla de los resultados, preferiblemente en el formato "CSV", es decir, como "valores separados por comas". Esta tabla se crea y guarda automáticamente a partir de la lista de códigos de frecuencia. Todos los datos aparecen aquí fila por fila y separados por comas. En la primera columna sólo hay letras como etiquetas de las variables para exportarlas a programas estadísticos, por ejemplo, SPSS, R o el módulo AQUAD para el análisis exploratorio de datos.

11.5 Combinación de métodos con el ejemplo de un estudio sobre el humor

En la siguiente sección, se demostrará la combinación de métodos cualitativos y cuantitativos utilizando el ejemplo de un estudio de Leo Gürtler. Los métodos mixtos nos parecen interesantes sobre todo cuando se trata de analizar grandes conjuntos de datos (por ejemplo, cuestionarios con respuestas abiertas), en los que algunas preguntas parciales son de carácter cuantitativo y ambos enfoques pueden complementarse de forma integradora. En cada caso, la atención se centra en la conexión lógica del proceso de modelización y comprobación de la teoría.

En un estudio con cuestionario más amplio ($N = 363$, véase el ejemplo anterior sobre la codificación de secuencias) con estudiantes de escuelas secundarias y colegios, se encuestaron teorías subjetivas sobre el humor en el aula. Las preguntas centrales de la investigación fueron:

- ¿Cómo experimentan los alumnos el humor en su clase?
- ¿Pueden identificarse diferentes puntos de vista y líneas de actuación en este contexto y cómo están conectados?

La construcción del cuestionario, derivada de la teoría, desarrolló nueve áreas temáticas, cada una de las cuales conducía a una pregunta abierta. A estos nueve temas se les asignaron códigos de hablante en la codificación posterior con AQUAD para permitir la evaluación selectiva según estos temas. Algunos ejemplos de preguntas fueron:

- Pregunta 1: ¿Qué es para ti el humor?
- Pregunta 6: Quizá también conozcas experiencias negativas con el humor (en clase). Si las conoces, ¿qué pasaba allí?

Las respuestas se escribieron en archivos de texto individuales para cada persona y se creó un proyecto AQUAD a partir de ellas. A continuación, se codificaron las respuestas según su contenido y estructura. En este caso, el contenido se refiere a códigos como "humor exterior" o "atmósfera/clima social", mientras que la estructura se introdujo como aspecto de validez. Esto permitió comprobar si las preguntas se respondían de acuerdo con su significado. Las categorías de ejemplo aquí son "definición", una categoría estructural que es de esperar en la pregunta 1 – la pregunta sobre la definición subjetiva del humor – y no, por ejemplo, en la pregunta sobre las secuencias de acción de las secuencias didácticas humorísticas. La pregunta sobre las secuencias de acción dio lugar a categorías como "efecto(s) positivo(s)" o "causa(s) negativa(s)".

A este primer proceso de codificación le siguió un segundo en el que los códigos de contenido se combinaron en metacódigos con la ayuda de un análisis de contenido cualitativo. Esto redujo el número de códigos a un nivel manejable ($N = 105$ metacódigos de contenido, $N = 22$ códigos estructurales). Los códigos estructurales eran menos numerosos y para ello no era necesario un análisis de contenido cualitativo.

11.5.1 Análisis de varianza univariante de las frecuencias de los códigos para generar hipótesis de vinculación

A partir de los memos anotados durante la codificación, los metacódigos resultantes y a través de la impresión subjetiva de los investigadores, surgieron varias áreas de hipótesis que podían desglosarse en los nueve temas. Además, se observa que las chicas parecen producir muchos más textos que los chicos y los alumnos del Gymnasium más que los de la Realschule (aunque esto último también podría deberse a una diferencia de edad en la muestra). Esto proporcionó una primera razón para el uso de la metodología cuantitativa. Se realizó un análisis de varianza univariante

de dos factores con los factores género y tipo de escuela y la siguiente codificación de contraste a priori (contrastes de Helmert según Fox 2002, p. 142 y ss.):

	[1]	[2]	[3]
Escuela de gramática masculina (mG)	- 0,5	+ 0,5	- 0,5
Realschule masculino (mR)	- 0,5	- 0,5	+ 0,5
Escuela de gramática femenina (wG)	+ 0,5	+ 0,5	+ 0,5
Realschule femenino (wR)	+ 0,5	- 0,5	- 0,5

Así, se han realizado las siguientes comparaciones:

[1]: mG & mR versus wG & wR	-> prueba de género.
[2]: mG & wG frente a mR & wR	-> prueba del tipo de escuela
[3]: mG & wR versus mR & wG	-> prueba de interacción género x tipo de escuela.

El modelo lineal resultante fue altamente significativo, al igual que los factores género y tipo de escuela. La interacción género x tipo de escuela no mostró significación, aunque hubo una ligera tendencia a la interacción, como sugiere un gráfico. Por lo tanto, parecía justificado y también necesario realizar más análisis cualitativos por separado para los grupos mG, mR, wR y wG.

Para poder reconstruir adecuadamente la comprensión subjetiva de los alumnos, se produjeron hipótesis de enlace en gran número. Éstas se derivaron de consideraciones teóricas formadas por la revisión del material de datos, las categorías de codificación resultantes y los memos elaborados. En general, estas hipótesis de enlace se limitaron a los temas y se ubicaron dentro de cada pregunta. Sin embargo, algunas excepciones también se formularon a través de las preguntas y, por lo tanto, dieron lugar a comparaciones de hablantes, ya que a las preguntas se les asignaron códigos de hablantes como se mencionó anteriormente. A continuación, algunos ejemplos para ilustrar este procedimiento:

Nombre del código de la secuencia

FR_HYPO_2-0-26: Y /\$M2 Límites del humor
 Y MX_causa negativa (-)
 Y MX_justificación/Contexto de razonamiento

Aquí comprobamos si las causas negativas se mencionaban en la pregunta 2, "Límites del humor", y se justificaban, es decir, no sólo se expresaban sino que los alumnos las justificaban en términos de argumento. Tenga en cuenta que hemos incluido la pregunta 2 (está disponible como código de hablante) en la codificación de la secuencia. A través de este pequeño truco, siempre sabremos después que un resultado positivo de esta hipótesis de secuencia sólo puede localizarse en la pregunta 2. Formular hipótesis de vinculación de esta manera puede abrir muy fácilmente nuevas perspectivas. Por lo tanto, recomendamos incluir los códigos de los hablantes en las hipótesis de enlace para las hipótesis que se refieren a áreas limitadas, como las preguntas o los actos de habla. Para encontrar todos aquellos pasajes de texto en los que se produce la misma secuencia "causa negativa" y "justificación", pero que no están relacionados con la pregunta 2, sólo hay que poner el primer código al final y enlazarlo con "NO". Recuerde que los enlaces "NO" no deben estar al principio de una hipótesis de enlace. El código de la secuencia tendrá entonces este aspecto:

Nombre del código de la secuencia

FR_HYPO_X-X-XX: Y MX_causa negativa (-)
 Y MX_justificación
 NO /\$M2 Límites del humor

Siguiendo este procedimiento se obtuvo un gran número de hipótesis de enlace, que se guardaron por separado en archivos individuales y luego se leyeron y probaron en el AQUAD. En este proceso, un ordenador rápido es muy útil, ya que esta prueba de hipótesis es un proceso computacionalmente intensivo. Aquad procede de tal manera que sólo se incluyen las hipótesis que tienen al menos un acierto. Esto ahorra la búsqueda posterior de resultados que no sean iguales a CERO.

La siguiente parte cuantitativa del análisis se ocupó de la pregunta anterior sobre la búsqueda de secuencias típicas en las respuestas a través de las personas de los respectivos grupos (mR, mG, wR, wG) o a través de todas las personas independientemente de su pertenencia al grupo. Los supuestos básicos de los procedimientos de análisis de conglomerados se utilizaron para encontrar la prototipicidad en las estructuras de datos. A diferencia del tipo ideal que no existe realmente en el sentido de Max Weber (1988), el prototipo es un dato o incluso varios datos (ya sea persona(s), codificación(es) u otras unidades de análisis) que se puede encontrar realmente de forma empírica y que se produjo realmente. El prototipo se define como el dato o los datos que tienen la menor distancia a todos los demás datos.

Este procedimiento pertenece al grupo de procedimientos analíticos de cluster del análisis exploratorio de datos. Suelen funcionar en dos pasos: En el primer paso, se calculan las distancias de cada fecha con todas las demás; en el segundo paso, se forman los clusters posteriores a partir de estas distancias según diferentes criterios. El prototipo puede determinarse directamente a partir del primer paso. Según la definición anterior, es el conjunto de datos con la menor distancia si se suman todas las distancias por separado para cada fecha. Este prototipo puede verse ahora como el centro o la línea central de argumentación de los datos disponibles.

Los análisis posteriores vuelven a entrar en el campo del análisis cualitativo, por lo que vuelven a ser "mixtos". De este modo, mediante el conocimiento de los prototipos y sus "opuestos" (es decir, los datos con las máximas distancias a todos los demás conjuntos de datos, determinados por el máximo de las distancias sumadas por fecha), se pueden conectar ahora los análisis cualitativos específicos, sin duda también con los datos brutos. En lo sucesivo, llamaremos a los casos "opuestos" "valores atípicos". Aquí, las preguntas cualitativas resultantes son, por ejemplo:

- ¿Qué es lo que distingue a los prototipos de los valores atípicos?
- ¿De qué códigos se trata en cada caso?
- ¿Qué es especial, qué es sistemático, cuando se comparan prototipos y valores atípicos?

Aplicado a nuestro ejemplo, obtenemos afirmaciones sobre cuáles de las hipótesis de secuencia teóricamente preformuladas son las que representan nuestros datos (las afirmaciones de los alumnos sobre el humor) particularmente bien (entonces son afirmaciones prototípicas) o particularmente mal (entonces son valores atípicos). Una pista nos parece especialmente importante:

En la siguiente interpretación, no hay que centrarse demasiado en los prototipos que representan especialmente bien los datos globales. Considere los valores atípicos como elementos igualmente importantes. Las respuestas interesantes a sus preguntas de investigación sólo las encontrará en la comparación y en la tensión entre los patrones de respuesta típicos y atípicos.

Ahora nos gustaría darle algunas instrucciones prácticas para el procedimiento que acabamos de describir.

Análisis de datos: De AQUAD a la matriz de distancias

Como hemos explicado anteriormente, necesitamos transformar nuestras codificaciones en una matriz de distancia que indique qué codificación (nuestras hipótesis de secuencia) está más o menos alejada de todas las demás. Para ello, utilizamos las frecuencias de las hipótesis de secuencia por persona. Para ello, creamos tablas y hacemos que Aquad cuente las frecuencias (es decir, los códigos). En otras palabras, ¿qué código de secuencia ocurre con qué persona y con qué frecuencia? El resultado de nuestro recuento de códigos se exporta en formato .csv, que puede transferirse fácilmente a un programa de hoja de cálculo o leerse con el lenguaje de programación estadística "R" (2004). Nuestra tabla debería tener ahora este aspecto:

Rúbrica (línea 1):	Indica el nombre de la variable o la numeración.
Columnas:	Personas
Filas:	Hipótesis de secuencia
Celdas:	Frecuencias de las hipótesis de secuencia

Si, por la razón que sea, sus filas y columnas están invertidas, puede cambiar esto fácilmente en R transponiendo la matriz.

Evaluación estadística con "R"

"R" (2004) es un lenguaje de programación especialmente diseñado para las evaluaciones estadísticas. Es gratuito y puede obtenerse en Internet (<http://cran.r-project.org/mirrors.html>). Allí también encontrará información, introducciones y referencias bibliográficas sobre el trabajo con R. Lamentablemente, no podemos reproducirlas aquí. Una buena guía para la creación de matrices de distancia puede encontrarse en Handl (2002, 83 y ss.). No obstante, conviene introducir algunos comandos:

```
table <- read.csv ("filename_our_frequency_table.csv", header=TRUE)
# leer la tabla de frecuencias en formato .csv y asignarla a la variable "tabla".
# la tabla .csv tiene cabeceras (!), si no, hay que poner "header = FALSE"
tabla
# para la comprobación, salida simple de la tabla de lectura a la pantalla
```

Ahora siguen las instrucciones de Handl (2002). Tenga en cuenta que normaliza la matriz (nuestra tabla) antes de calcular las distancias para compensar las diferentes dispersiones causadas por las personas (Handl, 2002, p. 97). Se necesita la matriz de distancia completa para determinar los prototipos (Handl, 2002, p. 96 y ss. o p. 437 resume un programa para esta tarea). Básicamente, se espera que las filas contengan las variables cuyas distancias entre sí se calculan. En nuestro ejemplo, se trata de las hipótesis de secuencia. Por lo tanto, las columnas contienen la base de datos para la matriz de distancia. En nuestro ejemplo, son las personas. Las frecuencias de cada código de secuencia por alumno pueden encontrarse en cada celda. Si tienes filas y columnas invertidas, puedes cambiar esto usando la matriz transpuesta:

```
tabla <- t(tabla)
# Intercambiar filas y columnas transponiendo la matriz "tabla".
```

Utilice la distancia "euclidiana" para la matriz de distancias, a menos que tenga presupuestos teóricamente justificados que sugieran una medida de distancia diferente. Obtenemos el prototipo de la matriz de distancia sumando todos los valores por fila. Obsérvese que en la matriz de distancia resultante, el orden de las filas y las columnas es el mismo, ya que dentro de la matriz de distancia los valores de las celdas indican la distancia de las respectivas variables de columna frente a las de fila. En la diagonal están los valores NULL, ya que es la distancia de cada valor a sí mismo, es decir, la identidad. Un valor de distancia de CERO significa máxima similitud, mientras que un valor de UNO significa máxima disimilitud.

Así que creamos el vector prototipo:

```
prototipo <- aplicar(matriz de distancia, 1, suma)
# Sumar todos los valores de la matriz de distancia con el nombre "distanzmatrix" por
fila y crear el vector prototipo con el nombre "prototype".
```

Guardamos este vector junto con la matriz de distancia en formato .csv.

```
write.table(cbind(prototype,distanzmatrix), file = "prototypetable.csv", sep = ",")
# Guarda el vector prototipo junto con la matriz de distancia bajo el nombre
"Prototipoentabelle.csv".
```

Ahora puede abrir la tabla resultante en el programa de hoja de cálculo de su elección. La columna de la izquierda designa el vector del prototipo. Ordena la tabla completa según este vector en orden ascendente. Esto le dará una secuencia de códigos de secuencia, donde los primeros valores (superiores) en el vector de prototipos tienen la mayor prototipicidad y los últimos (inferiores) tienen la menor prototipicidad - y por lo tanto son valores atípicos. Ahora compare a qué declaraciones de contenido corresponden los respectivos códigos de secuencia. En nuestro ejemplo, el resultado fue el siguiente para la pregunta 1 ("¿Qué es el humor?") en todos los estudiantes sin tener en cuenta la pertenencia a un grupo. El número que aparece después de "lugar" indica la clasificación de la prototipicidad: Aquí, el rango número 1 es prototípico, el rango número 2 ya lo es algo menos y los valores más altos son aún menos prototípicos.

```
Prototipo de lugar número 1: Y M_outward facing H Y M_laugh. funny. Diversión
Prototipo lugar 1: Y M_laugh. divertido. Diversión Y M_exteriormente orientado a H
Prototipo de posición 2: Y M_humor clásico Y M_reírse. divertido. Diversión
Posición del prototipo 2: Y M_classic humour Y M_borderline/ closeness to (-)
```

Se observa que los dos primeros códigos de secuencia son idénticos, pero sus elementos aparecen en distinto orden. Como ambos tienen el mismo valor de prototipicidad, pueden considerarse equivalentes. Por lo tanto, esta muestra está representada principalmente por el hecho de que para los estudiantes, independientemente de su afiliación a un grupo, el humor significa algo asociado con reírse, ser gracioso y divertirse, pero al mismo tiempo tiene un componente orientado hacia el exterior. Esto podría ser un indicio de que el humor es un asunto social e interactivo y no uno que tiene lugar solo en casa frente al televisor. Esto también podría considerarse una prueba de que los alumnos tenían en mente el contexto de la clase cuando respondían a la pregunta y, por tanto, esto podría respaldar la validez de las afirmaciones.

Sin embargo, en el caso del rango 2, se observa que además del humor clásico como "hacer bromas", también se asocia con reírse, ser gracioso y divertirse, y esto es después y no antes del humor clásico - ¡recuerda, el orden de codificación es importante aquí! Sin embargo, este código también va seguido de la categoría "límite/cercanía a (-)" y con el mismo valor de prototipicidad. Esto indica que, teniendo en cuenta que la atención se centra en el contexto de la clase, el humor puede tener no sólo aspectos positivos, sino también fuertemente negativos o ser experimentado negativamente. Y esto ya se observa a lo largo de la definición. Debido al alto valor de prototipicidad, hay que suponer que también es un punto de argumentación importante para los alumnos y los representa muy bien. Ahora podríamos

buscar las frecuencias absolutas de las hipótesis de secuencia mencionadas o sus relativas (normalizadas a los tamaños de los grupos) para obtener más información sobre la distribución de estas codificaciones. A continuación, podemos dejar que AQUAD nos muestre pasajes ejemplares del material de datos en bruto directamente en el texto para recoger las declaraciones originales como prueba y poder compararlas.

Con este procedimiento, se puede examinar una gran variedad de hipótesis, incluso con muestras muy grandes, y considerarlas en un contexto significativo. En particular, las hipótesis de secuencia pueden interpretarse directamente, ya que la secuencia fija deja claro cómo se conectan los códigos.

Como sugerencia adicional, nos gustaría llamar su atención sobre el trabajo de Oldenburger (1981, 153 ss.). Allí se describe cómo se puede utilizar una matriz de proximidad (una matriz de distancia es una matriz de este tipo) para llevar a cabo una dicotomización de base empírica, en conceptos que se conectan entre sí y conceptos que no se conectan entre sí. De este modo, no sólo se pueden identificar prototipos, sino que también se pueden formar clusters que, a diferencia de los métodos "clásicos" de agrupación jerárquica, pueden entrar en conexión entre sí y no están estrictamente separados. Afortunadamente, esta división se basa directamente en el empirismo y es compatible con el procedimiento según la teoría fundamentada (Glasser y Strauss, 1965). La salida es una matriz cero-uno en la que se puede leer para cada código con quién está relacionado (valor = UNO) o no (valor = CERO). La suma sobre las filas vuelve a dar como resultado el vector prototipo, pero aquí el máximo es el prototipo y no el mínimo, ya que un UNO significa "conexión existente" y, por lo tanto, distancia pequeña. Los CERO significan "sin conexión" y, por tanto, se resumen en una baja prototípica. La diagonal contiene mejor los CERO, ya que equivale a las varianzas y, por tanto, no contiene información redundante. También hay un programa en "R" para esto, que hace esta división de la matriz de distancia, Oldenburger lo llama "corte óptimo". Se llama "OptSchnPr" y está incluido en el módulo "Proxim" (Oldenburger, 2004) (véase el enlace web). El autor de AQUAD puede obtener un programa adicional de "R" que realiza la entrada/salida con tablas .csv y que puede ser adaptado por usted.

11.5.2 Formación de tipos por análisis de implicación sobre una base cuantitativa

Las hipótesis y conjeturas sobre la formación de elementos implicados se formulan generalmente de forma cualitativa y se basan en los "sospechosos habituales", a saber, los memos, las comparaciones inter e intraindividuales sobre el material de datos original, las hipótesis de vinculación y las secuencias de codificación. Sin embargo, a veces las diferencias empíricas son muy finas o la complejidad de los datos es muy elevada. Esto puede deberse a la novedad de un tema que apenas se ha estudiado empíricamente antes, o a otras razones. Esta situación hace necesarios enfoques alternativos para la generación de hipótesis que van más allá de los descriptos hasta ahora. Nos gustaría presentar un enfoque que utiliza sistemáticamente el análisis de los aspectos cuantitativos de los datos para llegar a hipótesis más claras sobre las relaciones que pueden conducir eventualmente a una tipología.

Las herramientas estadísticas utilizadas son muy sencillas. En primer lugar, se cuentan todos los códigos potencialmente relacionados (o metacódigos). A continuación, se lleva a cabo una transformación en z de la frecuencia por codificación en todas las personas (o en las entrevistas o las unidades temáticas que se van a analizar en sí). A continuación, los valores de todas las codificaciones/metacódigos se correlacionan entre sí. A partir de los valores extremos (por ejemplo, un valor z fuera de ± 2 desviaciones estándar) en los valores z de una persona investigada (o de una entrevista), se seleccionan las codificaciones correlacionadas con este código tan estadísticamente (altamente) significativas ($p \leq 0,05$) como sea posible y se determina un criterio para el análisis implícito en términos de contenido. Como se ha señalado anteriormente, no tenemos ante nosotros un diseño de investigación en el que las variables independientes y dependientes obliguen prácticamente a ciertas configuraciones. Más bien,

podemos formular y probar configuraciones plausibles, verosímiles, pero inicialmente arbitrarias, con diferentes predictores y criterios.

En un paso más, estos resultados deben ser evaluados (cualitativamente) para ver si las combinaciones de códigos encontradas de esta manera parecen ser relevantes en términos de contenido y para promover el conocimiento. En este caso, se realiza el análisis implícito en la selección realizada, si es posible para el criterio = VERDADERO y el criterio = FALSO (!). El objetivo de esta "doble contabilidad" es sensibilizar al investigador no sólo para que se fije en los resultados positivos, sino también para que saque conclusiones de las configuraciones de codificación negativas.

Esta secuencia de pasos descrita se repite hasta que una imagen coherente y justificable de las posibilidades de ordenación, por ejemplo en forma de tipología, no cambia significativamente. En el enfoque de la teoría fundamentada (Glaser & Strauss, 1965), se habla aquí de saturación, es decir, el material de datos ha empezado a revelar sus secretos y no cabe esperar que de los datos disponibles se encuentren más conocimientos sustanciales para poder responder a la pregunta inicialmente planteada de forma diferente a la anterior.

Sin embargo, una vez más, hay que llegar a un punto medio cuando se trata del número de implicados que hay que probar. Tampoco deben realizarse tantos análisis implicados en el famoso "enfoque de escopeta" de los análisis cuantitativos que sea necesario otro test de significación para separar el grano de la paja, ni deben basarse las argumentaciones más complejas sobre varias personas en un único criterio implicado. En el primer caso, surge el problema de que con un número creciente de análisis de implicados, también aumenta la probabilidad de que surjan combinaciones aparentemente significativas "por casualidad". Sin embargo, éstas también nublan la visión de las correlaciones reales, simplemente porque parecen plausibles y se adoptan "a ciegas". En el segundo caso se plantea el problema de que las interpretaciones basadas en muy pocos implicados pretenden un grado de generalización que los análisis realizados previamente no permiten en realidad debido a la elevada reducción de datos. Por lo tanto, las interpretaciones no están justificadas de forma exhaustiva. Por lo tanto, parece que merece la pena identificar al menos una configuración de codificación significativa por persona (entrevista u otra unidad de análisis) que ponga de manifiesto lo esencial. Esto aumenta la posibilidad de que se encuentren correlaciones de las configuraciones de condiciones también en varias unidades de análisis. Y es precisamente la tensión entre las configuraciones de las condiciones ("¿Por qué es así con ésta, mientras que no es así con la otra?") lo que permite extraer conclusiones con apoyo empírico recurriendo al material de datos original. La interpretación no se basa entonces en el resultado de un análisis de implicados aleatorio, sino en los intentos de tipificación explícita. Y no hay que olvidar que los tipos ideales de Weber son prácticamente imposibles de encontrar en la realidad. Pero esto, a su vez, significa que no sólo buscamos tipos, sino que también nos interesan las conexiones y diferencias entre los distintos tipos.

Nos gustaría mostrarle un pequeño ejemplo que funciona con el procedimiento que acabamos de describir. En un estudio de Gürtler (2002-2004) sobre las teorías subjetivas de los profesores sobre el humor en contextos profesionales y privados, se descubrió que una persona conceptualizaba una idea muy idealizada y de color espiritual sobre el humor. Como se ha descrito anteriormente, se realizó una transformación z y se correlacionaron todos los metacódigos entre sí. El punto de partida fue el metacódigo que denota la actitud humorística ideal. Esto se tomó como criterio positivo de la configuración de la condición a probar. Ahora, todos los metacódigos que se correlacionan significativamente ($p \leq 0,05$) con este código se incluyeron sucesivamente como predictores en el modelo implícito y se probaron. Los resultados fueron los siguientes:

A = M_Conocimiento
 B = M_actitud_ideal_de_humor
 C = M_Comunicación_soc_IA

D = M_autoconocimiento_conciencia

E = M_Self_image_int_ext

F = M_Jokes_actv_evaluación

Criterio: Condición 2 / FALSO

Implicado/s: ACeF [6 9], acde [4 8 10].

Criterio: Condición 2 / VERDADERO

Implicado/s: ACDEf [1]

A partir de un análisis más profundo, fue necesario probar adicionalmente una configuración similar pero no idéntica:

A = M_ideal_Humour_Attitude

B = M_cognActv_sin_soc_IA

C = M_comunicación_soc_IA

D = M_people_unspecific

E = M_Conocimiento de sí mismo

F = M_Self_reference_reflection

G = M_Misconcepción_del_concepto

Criterio: Condición 1 / VERDADERO

Implicado/s: BCDEFG [1]

Nota; No hay casos para un criterio negativo.

De ambos resultados se desprenden tesis interesantes. Una tesis, por ejemplo, apoyó la suposición del investigador de que las actividades relacionadas con los chistes (humor) no tenían importancia en este caso concreto (ACDEf), pero que el autoconocimiento, la interacción social y la comunicación, así como la inclusión de la imagen de sí mismo, eran factores importantes. Esto fue confirmado.

Sin embargo, también existe la ganancia de conocimiento por el uso del criterio negativo (condición 2 = FALSO). Aquí se podría demostrar que en una siguiente configuración encontrada, el chiste juega un papel importante (casi como contraste) y esto lleva a que la actitud de humor ideal se convierta en FALSA (= criterio). En este caso (ACeF), el papel de la imagen de sí mismo desde una perspectiva intrapersonal y exterior pierde inmensamente en importancia.

En el segundo análisis con el modelo ligeramente modificado se prescindió del código "conciencia", pero se integraron como predictores las "actividades cognitivas no relacionadas con la interacción social" y la "autorreflexión". Por supuesto, un análisis detallado podría ir más allá en este punto.

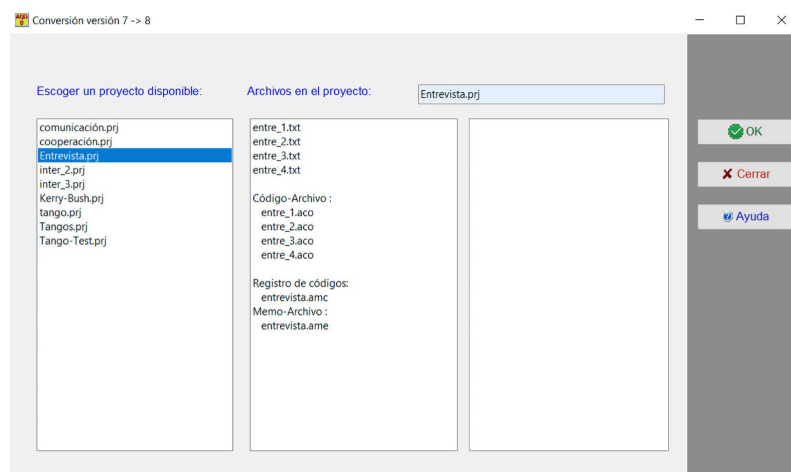
Capítulo 12: Conversión de códigos de formato AQUAD 7 a AQUAD 8

Inicialmente, AQUAD 8 no puede trabajar con los códigos de la versión anterior 7, ya que los códigos anteriores contenían no sólo 60 caracteres sino información adicional. Con esta función de conversión se consigue que los archivos antiguos se adapten a las condiciones de la versión 8.

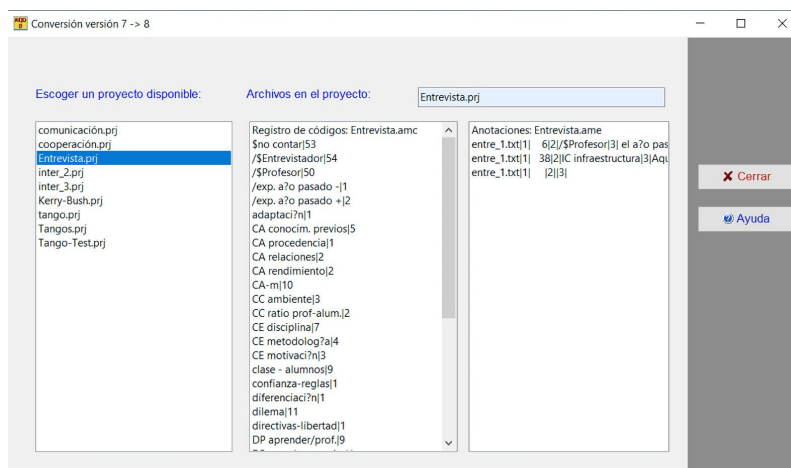
Sin embargo, la conversión sólo se ejecutará sin problemas si se han mantenido las convenciones vigentes cuando se trabaja con la versión antigua. En concreto, la rutina de conversión espera que

- los archivos de texto, sonido, vídeo o imagen de tu proyecto se han copiado en el directorio raíz (por ejemplo, C:\Aquad8_c_texto), los correspondientes archivos de proyecto, código y memo en el subdirectorio ..\cod de Aquad 8 (por ejemplo, C:\Aquad8_c_texto\cod);
- los nombres de los archivos de texto terminan en ".txt";

Cuando se llama a la rutina de conversión, aparece la ventana que se muestra a continuación. A la izquierda aparecen los nombres de todas las definiciones de proyectos disponibles. En la ventana de la derecha, seleccione el nombre del proyecto de la versión 7 que desea convertir, en este caso "*Entrevista.prj*". En el recuadro central, después de la selección, encontrará los nombres de los archivos (de texto), de los archivos de código, del registro de códigos y del archivo memo - si está disponible. Si a continuación hace clic en el botón "OK" del margen derecho, el resto se ejecuta automáticamente:



A continuación sigue:



Capítulo 13: Análisis estadísticos exploratorios

El módulo de estadística descriptiva-exploratoria según Tukey (1977) muestra que no es necesario calcular significancias para descubrir relaciones y diferencias significativas en los datos, sino que la creatividad y la adecuación de los casos, así como las transformaciones de los datos, son suficientes para derivar hipótesis sustanciales para investigaciones posteriores en combinación con diferentes fuentes de información. La idea rectora es siempre que, en lugar de realizar pruebas confirmatorias, es mejor examinar los datos por su contenido real para identificar estructuras y patrones. Proponemos prescindir de la estadística inferencial en el análisis de datos cualitativos y seguir sacando conclusiones razonables y sustanciales.

Programas de estadística permiten acceder a las siguientes funciones:

- Creación y edición de tablas de datos CSV
- Estadística exploratoria-descriptiva de datos de frecuencia (parámetros de distribución, correlaciones y gráficos de datos)
- Estadística exploratoria-clasificatoria de los datos de frecuencia (análisis de conglomerados, escalamiento multidimensional y cálculo de prototipos) y
- la representación de los datos de tablas CSV en diferentes diagramas (también basados en otros programas gratuitos para el análisis de tablas).

En el análisis exploratorio-estadístico, se utilizan métodos gráficos de visualización de datos y transformaciones (no) lineales de datos, en particular para destacar claramente las diferencias, correlaciones, tendencias y similares. Curiosamente, la tendencia a probar modelos complejos (por ejemplo, regresión / anova, HLMs / MLMs, ...) también apunta cada vez más al uso de métodos gráficos en lugar de basar los argumentos en los coeficientes solamente. Detrás del análisis exploratorio-estadístico se encuentran sobre todo transformaciones de datos y formaciones de subgrupos, así como diversas visualizaciones para poder ordenar, estructurar o simplemente mostrar los datos desde diferentes puntos de vista y según criterios de interés.

13.1 Modificación de las tablas de datos

Al determinar las frecuencias, por ejemplo, las frecuencias de códigos o palabras y el análisis de frecuencias de las tablas, Aquad 8 guarda los resultados en tablas CSV para su posterior uso en los módulos de análisis incorporados o en programas externos. Para muchos análisis, es útil o necesario editar primero estas tablas, por ejemplo, para eliminar columnas individuales (como la columna con la suma total de frecuencias en cada fila de la tabla), para combinar columnas, para transformar toda la tabla (intercambiar columnas y filas), etc.

13.2 Estadísticas exploratorias-descriptivas

Este módulo permite el análisis de las frecuencias de los códigos en un análisis cualitativo de texto, imagen, audio o vídeo, de las frecuencias de las palabras en un análisis cuantitativo de texto, la descripción de la distribución de las frecuencias, el cálculo de las correlaciones y los trazados bidimensionales o tridimensionales (es decir, las representaciones gráficas de las curvas y los datos) para el análisis gráfico exploratorio de los datos.

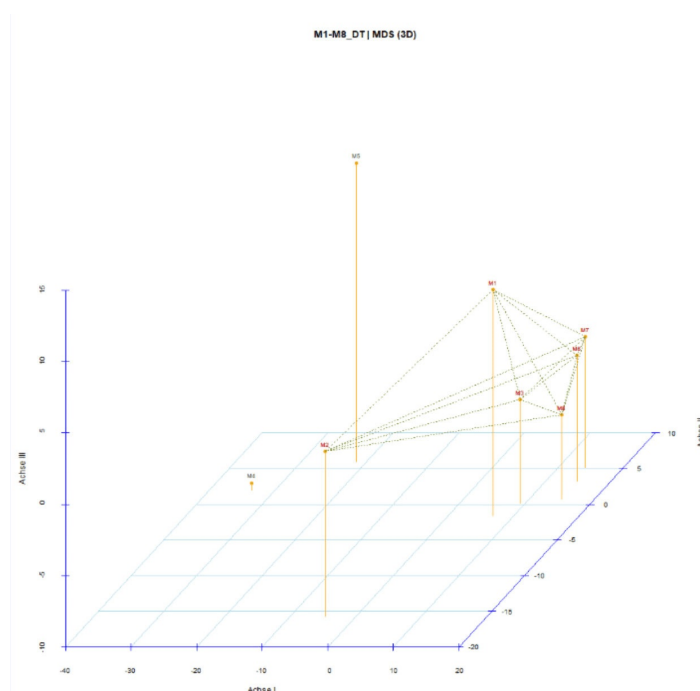
Importantes parámetros descriptivos de la distribución de datos en una tabla CSV son

- Número de casos (todos)
- Número de valores que faltan (NA)
- Número de observaciones (sin valores perdidos),
- Suma
- Media aritmética (MW)
- Media geométrica
- Media armónica
- Varianza (VAR)
- Desviación estándar (SD)
- Coeficiente de variación (CV)
- Mediana
- Desviación absoluta de la mediana (MAD)
- Mínimo
- Máximo
- Gama
- Desviación media
- Asimetría
- Kurtosis (exceso)
- 1er cuantil
- 3er cuantil
- Rango intercuartil (IQR)
- Error estándar de la media (S.E. MW)
- Media del intervalo de confianza inferior (LCL)
- Media del intervalo de confianza superior (UCL)

Otros análisis descriptivos proporcionan cálculos de correlación y la representación de los datos en gráficos, en los que se muestra la distribución de las variables individuales (columnas) de la tabla original en forma de diagramas de caja.

Si la evaluación clasificatoria tiene sentido, se pueden analizar con mayor precisión las tablas de frecuencias mediante análisis de conglomerados, escalamiento multidimensional y la determinación de prototipos.

Una escala tridimensional de los datos de una tabla CSV generada en AQUAD8 proporciona, entre otras cosas, el siguiente gráfico:



Sin duda, el escalamiento multidimensional visualiza las correlaciones entre las variables, aquí las frecuencias de los metacódigos seleccionados de un estudio de Gürtler, y sus agrupaciones de forma impresionante y puede ser muy informativo en cuanto a la pregunta de investigación específica. Pero, decisivamente, que este análisis tenga sentido depende de la pregunta de investigación.

Lo mismo ocurre con el análisis de los **prototipos**. Además de varias tablas, este script de R proporciona una visualización muy clara del corte óptimo a través de la matriz de proximidad de la tabla de salida. Si no se entienden los conceptos básicos y su significado, no se deben utilizar estos algoritmos de clasificación en plan "a ver qué sale". En la estadística inferencial clásica, se utilizaría el término "enfoque de escopeta" para este procedimiento ("algo será significativo").

13.3 Análisis comparativo cualitativo / minimización booleana

Después de la descripción detallada del método en el capítulo 10, a continuación se ofrece un breve comentario sobre el significado y la finalidad de la minimización booleana en el análisis de datos cualitativos. Especialmente en el estudio de la experiencia y el comportamiento social, no se debe esperar encontrar una única causa específica para un fenómeno determinado. Más bien, es importante descubrir las constelaciones de condiciones en las que se suele esperar un fenómeno crítico. Para ello, hay que comparar muchos casos en los que se produce este fenómeno. Esto va necesariamente acompañado de una cierta falta de claridad, que, sin embargo, no es "perjudicial" ni "obstruiva", sino que corresponde a la realidad. Ragin (1987) ha abierto un procedimiento explícitamente comparativo aplicando el álgebra de Boole a los datos cualitativos. El término análisis comparativo cualitativo (ACQ) se utiliza a menudo en la literatura. Dado que el objetivo del análisis es reducir las condiciones iniciales a configuraciones mínimas de condiciones, llamadas implicantes, utilizaremos el nombre de análisis de implicantes, que es sinónimo de ACQ.

El objetivo es explicar un conjunto definido de resultados con un conjunto mínimo de factores condicionales. Aplicado a los fenómenos sociales, el enfoque se utiliza para encontrar el conjunto mínimo de condiciones para que

un criterio sea lógicamente VERDADERO o FALSO en todos los casos individuales examinados (pueden ser personas, textos, pero también estudios enteros). Debido a su flexibilidad, el algoritmo también es adecuado para el meta-análisis.

Además de un algoritmo propio para la minimización booleana, que se describe detalladamente en el capítulo 10 y que está disponible en AQUAD desde la versión 4, hasta la versión Aquad 8.5 también se utilizó el paquete R QCA (Qualitative Comparative Analysis) para determinar implicantes. Sin embargo, dado que este paquete parece causar problemas en las versiones más recientes de R y que los componentes necesarios ya no están disponibles en otros contextos, hemos eliminado «aquad_eda.exe» del paquete AQUAD. El enlace a los algoritmos propietarios de AQUAD sigue estando disponible. Para obtener más detalles y ejemplos sobre el uso del módulo AQC, consulte el capítulo 10 más arriba.

Como se ha mencionado al principio, para la visualización de las distribuciones se dispone de todas las posibilidades de un programa de hoja de cálculo. En concreto, el módulo "scal" del software libre "LibreOffice" abre la tabla CSV, que se va a procesar, y ofrece las siguientes opciones de visualización, tanto bidimensional como tridimensional:

- Gráficos de barras
- Gráficos de barras
- Gráficos circulares
- Diagramas de área
- Gráficos de líneas
- Gráficos de dispersión (X,Y)
- Gráficos de burbujas
- Diagramas de red
- Diagramas del curso
- Columnas y líneas

Por supuesto, se pueden seleccionar o eliminar rangos de datos, cambiar los rótulos de los ejes, cambiar el color de los elementos de los diagramas y del fondo de los mismos, por enumerar sólo algunas opciones importantes.

El requisito previo para utilizar el módulo de diagramas es tener instalado el software gratuito "LibreOffice" en su ordenador.

Referencias

- Berelson, B. (1952). *Content analysis in communication research*. Glencoe, Ill.: The Free Press.
- Creswell, J. W. (2012). *Qualitative Inquiry and Research Design: Choosing Among Five Approaches*. Thousand Oaks, CA: Sage Publications.
- Eisner, E. W. (1981). On the differences between scientific and artistic approaches to qualitative research. *Educational Researcher*, 10(4), 5-9.
- Eisner, E. W. (1983). Anastasia might still be alive, but monarchy is dead. *Educational Researcher*, 12(5), 13-14/23-24.
- Erzberger, C., & Kelle, U. (2003). Making inferences in mixed methods: The rules of integration. In A. Tashakkori & C. Teddlie (Eds.), *Handbook of mixed methods in social and behavioral research* (pp. 457-488). Thousand Oaks: Sage.
- Fetterman, D. M. (1982). Ethnography in educational research: The dynamics of diffusion. *Educational Researcher*, 11(3), 17-22.
- Fetterman, D. M. (1988). Qualitative approaches to evaluating education. *Educational Researcher*, 17(8), 17-24.
- Firestone, W. A. (1987). Meaning in method: The rethoric of quantitative and qualitative research. *Educational Researcher*, 16(7), 16-21.
- Fischer, P. M. (1994). Inhaltsanalytische Auswertung von Verbaldaten. In G. L. Huber & H. Mandl (Hrsg.), *Verbale Daten* (2. Aufl.) (S. 179-196). Weinheim: Beltz.
- Fox, J. (2002). *An R and S-PLUS companion to applied regression*. Thousand Oaks: Sage .
- Fühlau, I. (1978). Untersucht die Inhaltsanalyse eigentlich Inhalte? Inhaltsanalyse und Bedeutung. *Publizistik*, 7-18.
- Fühlau, I. (1982). *Die Sprachlosigkeit der Inhaltsanalyse. Linguistische Bemerkungen zu einer sozialwissenschaftlichen Methode*. Tübingen: Gunter Narr.
- Galtung, Johan (1990). Theory formation in social research: A plea for pluralism. In E. Øyen (Ed.), *Comparative methodology: Theory and practice in international social research* (S.96--112). London: Sage.
- Gläser-Zikuda, M., Seidel, T., Rohlf, C., Gröschner, A., & Ziegelbauer, S. (Hrsg.) (2012). *Mixed methods in der empirischen Bildungsforschung*. Münster: Waxmann.
- Glaser, B. G. (1978). *Theoretical sensitivity*. Mill Valley, CA: Sociology Press.
- Glaser, B. G. (1992). *Emergence vs. forcing. Basics of grounded theory analysis*. Mill Valley, CA: Sociology Press.
- Glaser, B. G. & Strauss, A. L. (1965). *The discovery of grounded theory. Strategies for qualitative research*. Chicago: Aldine Publishing Company.

Glaser, B. G., & Strauss, A. L. (1979). Die Entdeckung gegenstandsbezogener Theorie: Eine Grundstrategie qualitativer Sozialforschung. In C. Hopf & E. Weingarten (Hrsg.), *Qualitative Sozialforschung* (S. 91-111). Stuttgart: Klett-Cotta.

Glass, G. V., McGaw, B., & Smith, M. L. (1981). *Meta-analysis in social research*. Beverly Hills: Sage.

Greene, J. C., Caracelli, V. J., & Graham, W. F. (1989). Toward a conceptual framework for mixed-methods evaluation designs. *Educational Evaluation and Policy Analysis*, 11 (3), 255-274.

Günther, U. L. (1987). Sprachstil, Denkstil und Problemlöseverhalten. Inhaltsanalytische Untersuchungen über Dogmatismus und Abstraktheit. In P. Vorderer & N. Groeben (Hrsg.), *Textanalyse als Kognitionskritik?* (S. 22-45). Tübingen: Narr.

Gürtler, L. & Huber, G. L. (2012). Triangulation. Vergleiche und Schlussfolgerungen auf der Ebene der Datenanalyse. In M. Gläser-Zikuda, T. Seidel, C. Rohlf, A. Gröschner & S. Ziegelbauer (Hrsg.), *Mixed methods in der empirischen Bildungsforschung* (pp.37-50). Münster: Waxmann.

Gürtler, L., Studer, U. M., & Scholz, G. (2010): *Tiefensystemik, Bd. 1. Lebenspraxis und Theorie: Wege aus Süchtigkeit finden*. Münster: Verlag Monsenstein und Vannerdat, MV-Wissenschaft.

Härtling, P. (o.J.): *Das war der Hirbel*. Stuttgart: Klett-Cotta.

Handl, A. (2002). *Multivariate Analysemethoden. Theorie und Praxis multivariater Verfahren unter besonderer Berücksichtigung von S-PLUS*. Berlin: Springer.

Held, J. (1994). *Praxisorientierte Jugendforschung. Theoretische Grundlagen, methodische Ansätze, exemplarische Projekte*. Hamburg: Argument.

Held, J., Horn, H.-W. & Marvakis, A. (1996). *Gespaltene Jugend. Politische Orientierungen jugendlicher ArbeitnehmerInnen*. Opladen: Leske + Budrich.

Held, J., & Marvakis, A. (1992). Empirische Jugendforschung und ihr Verhältnis zur politischen Bildung. In K.-H. Braun & K. Wetzel (Hrsg.). *Lernwidersprüche und pädagogisches Handeln*. Bericht von der 6. internationalen Ferienuniversität Kritische Psychologie, 24. bis 29. Februar 1992 in Wien (S. 243-256). Marburg: Verlag Arbeit & Gesellschaft.

Hildenbrand, B. (2006): *Einführung in die Genogrammarbeit* [Introduction to the work with genograms]. Heidelberg: Carl Auer.

Howe, K. R. (1985). Two dogmas of educational research. *Educational Researcher*, 14 (8), 10-18.

Howe, K. R. (1988). Against the quantitative-qualitative incompatibility thesis or dogmas die hard. *Educational Researcher*, 17 (8), 10-16.

Huber, A. A. (2007). How to add qualitative profundity to quantitative findings in a study on cooperative learning. In Ph. Mayring, G.L. Huber, L. Gürtler, & M. Kiegelmann (Eds.) *Mixed methodology in psychological research* (pp. 179-190). Rotterdam: Sense Publishers.

- Huber, G. L. (Ed.) (1992). *Qualitative Analyse. Computereinsatz in der Sozialforschung*. München: Oldenbourg.
- Huber, G. L. (1992). Qualitative Analyse mit Computerunterstützung. In G. L. Huber (Hrsg.), *Qualitative Analyse. Computereinsatz in der Sozialforschung* (S. 115-176). München: Oldenbourg.
- Huber, G. L. (2001). *Typenbildung in der qualitativen Analyse am Beispiel der Lehrerforschung*. Beitrag zur Tagung der Fachgruppe Pädagogische Psychologie der Deutschen Gesellschaft für Psychologie, Landau 2001.
- Huber, G. L. & Kenntner, S. (1988). *Struktur biographischer Selbst-Schemata*. Bericht an die DFG. Tübingen: Arbeitsbereich Pädagogische Psychologie am Institut für Erziehungswissenschaft I der Universität Tübingen.
- Huber, G. L., & Marcelo García, C. (1992). Voices of beginning teachers: Computer-assisted listening to their common experiences. In M. Schratz (Ed.), *Qualitative voices in educational research*. (S. 139-156). London: Falmer.
- Huber, G. L., & Roth, J. W. H. (2004). *Research and intervention by interviews and learning diaries in inservice teacher training*. Paper presented at the workshop on "Mixed Methods in Psychological Research" (organized by SIG #17, EARLI and the Center for Qualitative Psychology, University of Tübingen) at Freudenstadt, Germany, Oct. 21-24, 2004.
- Huberman, M. (1987). How well does educational research really travel? *Educational Researcher*, 16 (1), 5-13.
- Jackson, G. B. (1980). Methods for integrative reviews. *Review of Educational Research*, 50, 438-460.
- Jacob, E. (1988). Clarifying qualitative research: A focus on traditions. *Educational Researcher*, 17 (1), 16-24.
- Jordell, K. (1987). Structural and personal influence in the socialization of beginning teachers. *Teaching and Teacher Education*, 3, 165-177.
- Kelle, U. (Ed.) (1995). *Computer-aided qualitative data analysis. Theory, methods, and practice*. London: Sage.
- Kiegelmann, M., Held, J., Huber, G. L. & Ertel, I. (2000). Das Zentrum für Qualitative Psychologie an der Universität Tübingen [24 Absätze]. *Forum Qualitative Sozialforschung; Forum: Qualitative Social Research [On-line Journal]*, 1/(2). Disponible en:
<http://www.qualitative-research.net/fqs-texte/2-00/2-00kiegelmannetal-d.htm>
- Kracauer, S. (1952). The challenge of qualitative content analysis. *Public Opinion Quarterly*, 16, 631- 642.
- Lisch, R. & Kriz, J. (1978). *Grundlagen und Modelle der Inhaltsanalyse*. Reinbek bei Hamburg: Rowohlt.
- Marcelo García, C. (1991). *El primer año de enseñanza*. Sevilla: Grupo de Investigación Didáctica de la Universidad de Sevilla.
- Marcelo García, C. (1992). *Desarrollo profesional e iniciación a la enseñanza. Estudio de caso de un programa de formación para profesores principiantes*. Resolución de 30 de sept. de 1992, de la Universidad de Sevilla.
- Mathison, S. (1988). Why triangulate? *Educational Researcher*, 17 (2), 13-17.
- Mayring, Ph. (1988). *Qualitative Inhaltsanalyse*. Weinheim: Deutscher Studien Verlag.

Mayring, Ph. (2001, Februar). Kombination und Integration qualitativer und quantitativer Analyse [31 Absätze]. *Forum Qualitative Sozialforschung / Forum: Qualitative Social Research* [On-line Journal], 2(1). Verfügbar über: <http://qualitative-research.net/fqs/fqs.htm>

Mayring, Ph. (2007). Mixing qualitative and quantitative methods. In Ph. Mayring, G.L. Huber, L. Gürtler, & M. Kiegelmann (Eds.) *Mixed methodology in psychological research* (pp. 27-36). Rotterdam: Sense Publishers.

Mayring, Ph., Huber, G.L., Gürtler, L., & Kiegelmann, M. (Hrsg.) (2007). *Mixed methodology in psychological research*. Rotterdam: Sense Publishers.

Maxwell, J. (2002). Realism and the roles of the researcher in qualitative psychology. In M. Kiegelmann (Ed.), *The role of the researcher in qualitative psychology* (pp. 11-30). Tübingen: Ingeborg Huber Verlag.

Medina, A., Feliz, T., Domínguez, M.-C., & Pérez, R. (2002). The methodological complementariness of biograms, in-depth interviews, and discussion groups. In M. Kiegelmann (Ed.), *The role of the researcher in qualitative psychology* (pp. 169-184). Tübingen: Ingeborg Huber Verlag.

Miles, M. B. (1985). Qualitative data as an attractive nuisance: The problem of analysis. In J. Van Maanen (Ed.), *Qualitative methodology*. (5. Aufl.). (S. 117-134). Beverly Hills: Sage.

Miles, M.B. & Huberman, A.M. (1984). *Qualitative data analysis: A sourcebook of new methods*. Beverly Hills, CA: Sage.

Miles, M. B., & Huberman, A. M. (1994). *Qualitative data analysis. An expanded sourcebook*. (2nd edition). Thousand Oaks, CA: Sage.

Oldenburger, H. (1981). *Methodenheuristische Überlegungen und Untersuchungen zur "Erhebung" und Repräsentation kognitiver Strukturen*. Göttingen: Dissertation.

Oldenburger, H. (2004). *Proxim: Repräsentation von Proximitymatrizen - Beschreibung und Analyse*. <http://www.liteline.de/~holdenb/fst/nwz/R-PHP/Proxim.R>

Oevermann, U. (1996). *Konzeptualisierung von Anwendungsmöglichkeiten und praktischen Arbeitsfeldern der objektiven Hermeneutik*. Frankfurt am Main: Goethe-Universität.

Oevermann, U. (2002). *Klinische Soziologie auf der Basis der Methodologie der objektiven Hermeneutik*. Frankfurt am Main: Goethe-Universität. Disponible en: <http://publikationen.ub.uni-frankfurt.de/frontdoor/index/index/docId/4958> [Zugriff am 22.01.2012].

Oevermann, U., Allert, T., Konau, E., & Krambeck, J. (1979). Die Methodologie einer "objektiven Hermeneutik" und ihre allgemeine forschungslogische Bedeutung in den Sozialwissenschaften (pp. 352-434). In H.-G. Soeffner (Ed.), *Interpretative Verfahren in den Sozial- und Textwissenschaften*. Stuttgart: Metzler.

Phillips, D. (1983). After the wake: Post-positivistic educational thought. *Educational Researcher*, 12(5), 4-12.

Popko, E. S. (1980). *Key-word-in-context bibliographic indexing. Release 4.0 users' manual*. Cambridge, Mass.: Harvard University, Laboratory for Computer Graphics and Spatial Analysis.

Popper, K. (1934/2005). *Logik der Forschung* (11. Aufl.). Tübingen: Mohr Siebeck.

R - Development Core Team (2004). *R: A language and environment for statistical computing*. Vienna: Austria.
<http://www.r-project.org>

Ragin, C. C. (1987). *The comparative method. Moving beyond qualitative and quantitative strategies*. Berkeley: University of California Press.

Reichert, J. (2000). Zur Gültigkeit von Qualitativer Sozialforschung [76 Absätze]. *Forum Qualitative Sozialforschung / Forum: Qualitative Social Research [On-line Journal]*, 1(2). Abrufbar über: <http://qualitative-research.net/fqs/fqs-d/2-00inhalt-d.htm> [Zugriff: 22.01.2012].

Saldaña, J. (2016). *The coding manual for qualitative researchers*. London: Sage Publications.

Schatz, M. (Ed.) (1993). *Qualitative voices in educational research*. London: Falmer Press.

Schweizer, H. (2004). Preparations for the Redemption of the World: Distribution of Words and Modalities in Chapter I of Don Quixote. In M. Kieglmann & L. Gürtler (Eds.), *Research Questions and Matching Methods of Analysis* (pp. 71-108). Tübingen: Ingeborg Huber Verlag.

Shelly, A. L. (1986). The stages of qualitative data analysis. In A. L. Shelly, R. A. Archambault, C. Sutton & P. Tinto (Eds.), *The stages of qualitative research: Researcher thinking facilitated by logic programming*. CASE Technical Report No. 8607, Syracuse, NY: Syracuse University, The Center for Computer Applications and Software Engineering.

Shelly, A. L., & Sibert, E. E. (1985). *The QUALOG user's manual*. Syracuse, NY: Syracuse University, School of Computer and Information Science.

Shelly, A. L., & Sibert, E. E. (1992). Qualitative Analyse: Ein computerunterstützter zyklischer Prozeß. In G. L. Huber (Hrsg.), *Qualitative Analyse. Computereinsatz in der Sozialforschung* (S. 71-114). München: Oldenbourg.

Shulman, L. S. (1981). Disciplines of inquiry in education: An overview. *Educational Researcher*, 10, 5-12.

Smith, J. K. (1983). Quantitative versus qualitative research: An attempt to clarify the issue. *Educational Researcher*, 12(3), 6-13.

Smith, J. K. & Heshusius, L. (1986). Closing down the conversation: The end of the quantitative-qualitative debate among educational researchers. *Educational Researcher*, 15(1), 4-12.

Strauss, A., & Corbin, J. (1990). *Basics of qualitative research. Grounded theory procedures and techniques*. Newbury Park, CA: Sage.

Studer, U. M. (1988). *Verlangen, Süchtigkeit und Tiefensystemik*. Fallstudie des Suchttherapiezentrum für Drogenabhängige
 START AGAIN in Zürich zwischen 1992–1998. Bericht ans Bundesamt für Justiz.

Tashakkori, A., & Teddlie, C. (1998). *Mixed methodology: Combining qualitative and quantitative approaches* (Applied Social Research Methods, No. 46). Thousand Oaks: Sage.

Tashakkori, A., & Teddlie, C. (Eds.) (2003). *Handbook of mixed methods in social and behavioral research*. Thousand Oaks: Sage.

Tesch, R. (1990). *Qualitative research. Analysis types and software tools*. New York: Falmer.

Tesch, R. (1992). Verfahren der computerunterstützten qualitativen Analyse. In G. L. Huber (Hrsg.), *Qualitative Analyse. Computereinsatz in der Sozialforschung* (S. 43-70). München: Oldenbourg.

Thomas, W. I., & Znaniecki, F. (1918). *The Polish peasant in Europe and America*. Boston: Badger.

Tuthill, D. & Ashton, P. (1983). Improving educational research through the development of educational paradigms. *Educational Researcher*, 12(10), 6-14.

Van der Linden, J., Erkens, G., & Nieuwenhuysen, T. (1995). Gemeinsames Problemlösen in Gruppen. *Unterrichtswissenschaft*, 23, 301-315.

Villar, L. M. (Ed.) (1981). *Las prácticas de enseñanza*. Sevilla: ICE de la Universidad.

Villar, L. M., & Marcelo, C. (1992). Kombination qualitativer und quantitativer Methoden. In G. L. Huber (Hrsg.), *Qualitative Analyse. Computereinsatz in der Sozialforschung* (S. 177-218). München: Oldenbourg.

Weber, M. (1988). *Gesammelte Aufsätze zur Wissenschaftslehre*. Tübingen: Mohr.

Weber, R. P. (1985). *Basic content analysis*. Sage University Paper series on Quantitative Applications in the Social Sciences, series no. 07-049. Beverly Hills: Sage.

Wernet, A. (2009). *Einführung in die Interpretationstechnik der Objektiven Hermeneutik* (3rd ed.). Wiesbaden: VS Verlag für Sozialwissenschaften.

Yin, R. K. (1989). *Case study research. Design and methods*. Applied Social Research Methods Series, Vol. 5. Newbury Park, CA: Sage.

Zabalza, M. A. (1991). *Los diarios de clase. Documento para estudiar cualitativamente los dilemas prácticos de los profesores*. Barcelona: Promociones y Publicaciones Universitarias, S.A.

Sobre los autores

Prof. i.R. Dr. phil. Dr. h.c. Günter L. Huber (jubilado, profesor titular de Psicología de la Educación en el Instituto de Ciencias de la Educación de la Universidad de Tübinga). Su docencia, investigación y publicaciones incluyen temas de aprendizaje cooperativo, dificultades de enseñanza-aprendizaje, diferencias interindividuales y análisis de datos cualitativos. Fundador y desarrollador de AQUAD.

Contacto: info@aquad.de

Dr. rer. soc. Leo Gürtler (Dipl.-Psych., doctorado en Ciencias de la Educación) - muchos años de práctica en la investigación empírica (por ejemplo, humor, problemas de adicción, ciencia empírica de la educación) y profesor académico. Experiencia en el extranjero en el sector privado (sanidad) como Change Manager. Trabaja como coach sistémico y consultor científico. Contacto: info@guertler-consulting.de