

12. Singularidades a tener en cuenta en un proyecto de telemedicina rural

Virgilio Cane León¹, German Hirigoyen Emparanza² y
Pablo Javier Belzarena García³

En este capítulo se describe la clasificación de servicios de telemedicina en función del tiempo, y su implicación en la solución adoptada para zonas rurales. Se incluye la descripción de dos tipos de escenarios de telemedicina: simples y complejos. Además se resume en forma tabular las diferentes aplicaciones y equipamientos técnicos en telemedicina. Finalmente, se analizan un conjunto de características técnicas que inciden en la calidad de servicio (QoS) de las aplicaciones de telemedicina.

12.1. Servicios de telemedicina en tiempo real o en diferido

La multiplicidad de especialidades existentes en la medicina y las diversas maneras de adaptar o utilizar las TIC para apoyar su práctica hacen que existan varias clasificaciones de los servicios de telemedicina. En el Capítulo 7 se describen en detalle las clasificaciones basadas en el tipo de servicio y el tipo de estudio médico, mientras que aquellas basadas en el tiempo las trataremos a continuación.

La clasificación en el tiempo se refiere las circunstancias temporales en las que se realizan la intervención médica a distancia y la comunicación entre el proveedor del servicio y el cliente. Se tiene un servicio en tiempo real cuando proveedor y cliente coinciden en el momento de intercambio de información, como en los casos de una llamada telefónica, una conversación vía mensajería instantánea (*chat*), o una videoconferencia. Se tiene un servicio en diferido cuando la información es generada en un instante y recibida un tiempo después, como en el caso de un registro en una historia clínica que es diligenciado por el médico general y más tarde consultado por el especialista.

¹Instituto de Investigaciones en Ciencias de la Salud, Universidad Nacional de Asunción, Paraguay

²Fundación de Telemedicina (Fundatel), Argentina

³Universidad de la República, Uruguay

Uno de los factores que puede ser determinante para la selección del tipo de intervención, si en diferido o en tiempo real, es el ancho de banda de las comunicaciones, que depende del tipo de señal a transmitir, de su volumen y del tiempo de respuesta requerido por las aplicaciones. Una comparación que puede ilustrar la diferencia de requerimientos entre tiempo real y diferido es entre una sesión de teleconsulta, en tiempo real, que con condiciones de calidad mínima requiere entre 0,5 y 1 Mbps, mientras que el envío de imágenes estáticas o el monitoreo de signos vitales requieren anchos de banda inferiores [201].

12.1.1. Servicios de telemedicina en tiempo diferido (*store and forward*)

En esta situación el cliente de algún servicio de telemedicina no se halla en comunicación directa o no está en línea con el proveedor del servicio. A esta modalidad se la denomina como *store-and-forward* o de “tiempo diferido”. El proveedor del servicio almacena las solicitudes de telemedicina y en un período las procesa y luego las devuelve al cliente como resultado de su servicio.

Un caso para ilustrar el tiempo diferido sería el de los estudios estáticos de Medicina Nuclear, en los cuales el médico especialista recibe cierta cantidad de imágenes estáticas de centellografía paratiroides con el fin de evaluarlas desde su despacho, las examina y luego las devuelve con el diagnóstico sin haber tenido contacto directo con el paciente o con el técnico que ha realizado el estudio.

Con la modalidad de tiempo diferido los estudios a diagnosticar se almacenarán en el computador del médico especialista o en un servidor dedicado, y luego serán tratados uno a uno por el especialista, quien podrá enviar todos los resultados al mismo tiempo o hacerlo de a uno a medida que realiza el diagnóstico.

La gran mayoría de las aplicaciones de diagnóstico por telemedicina funcionan en tiempo diferido, ya que no requieren gran ancho de banda en las comunicaciones y por tanto su costo es bajo; además, no requieren coordinación temporal entre los especialistas y el paciente.

12.1.2. Servicios de telemedicina en tiempo real (*real time*)

El tiempo real hace mención al hecho de que el cliente y el proveedor se encuentran en contacto directo a través de un medio de comunicación. Se obtiene así una interacción personal, por ejemplo entre el médico y el paciente, que suele ser más eficaz que si se hiciera en tiempo diferido porque permite acceder a determinada información clínica que en algunos casos resulta imprescindible. Sin embargo, estos servicios requieren de anchos de banda superiores que generalmente son costosos, y además dependen de la disponibilidad simultánea de los actores [202].

Como ejemplos de este tipo de servicio se pueden mencionar la teleconsulta, la teleasistencia, la teleeducación interactiva, y las aplicaciones que presenten casos de urgencias que ameriten una transmisión en tiempo real.

Existen dos instrumentos fundamentales para la telemedicina en tiempo real: 1) la videoconferencia, que permite la comunicación bidireccional de audio y/o video, y 2) las aplicaciones interactivas, que son programas que permiten sincronizar dos aplicaciones remotas para que los actores de telemedicina puedan compartir la información al mismo tiempo. Por ejemplo, una aplicación interactiva de teleecografía mediante la cual un especialista puede mostrar detalles de la imagen a otro especialista en tiempo real y usar una función de filtro que será aplicada igualmente en la aplicación remota, de modo que los dos actores ven exactamente las mismas imágenes.

12.1.3. Escenarios de telemedicina

En los últimos años el desarrollo de Internet, la telefonía móvil y las nuevas redes de telecomunicaciones de banda ancha, han aumentado las capacidades tecnológicas que puede utilizar el sistema sanitario para ofrecer servicios de asistencia médica a distancia. Sin embargo, estas capacidades no están accesibles en forma equitativa para toda la población, por lo que es necesario diferenciar dos escenarios para la prestación de servicios de telemedicina, en función a su disponibilidad y ubicación geográfica: simples y complejos (integrales) [203].

La identificación de estos dos escenarios permite que se ponga en evidencia que para la puesta en marcha de servicios de telemedicina no solamente se debe evaluar la disponibilidad de especialistas y equipamientos médicos, sino además aspectos como el ancho de banda de las redes disponibles en el área de implementación, para así poder determinar correctamente la modalidad a utilizar.

Los escenarios simples son utilizados generalmente en telemedicina rural, donde los recursos son escasos y, por lo general, así existan varios servicios solo se puede utilizar uno a la vez debido a las limitaciones de ancho de banda de las redes existentes.

Los escenarios complejos, como las zonas urbanas, disponen de mayores recursos de comunicaciones que permiten muchos más usuarios conectados y la utilización de sistemas más sofisticados.

La diferencia entre ambos escenarios está determinada fundamentalmente por la disponibilidad del servicio y por el ancho de banda.

12.1.4. Aplicaciones y equipos por ámbito de utilización

Con el fin de ofrecer a la población una solución de telemedicina apropiada a sus necesidades, deben considerarse las distintas maneras de implementar las aplicaciones mediante las tecnologías disponibles. Para ello es importante contar con una base de conocimiento que a partir del ámbito o servicio específico deseado nos permita determinar qué aplicaciones son necesarias y qué equipos se pueden tener en cuenta para implementarlas.

Como referencia se muestra en la Tabla 12.1 las distintas aplicaciones de telemedicina correlacionadas con su utilización clínica y ámbitos de implementación, y en la Tabla 12.2 los equipos utilizados por estas aplicaciones.

Ámbito Aplicaciones	Rurales	Urbanas	Emergen- cias y desas- tres	Fronte- rizas	Trata- miento de pa- tologías específi- cas	Segunda opi- nión	Atención espe- cializa- da en salud	Remisión de pa- cientes
Evaluación inicial del estado de urgencia y transferencia	X	X	X		X			X
Tratamiento médico y post-quirúrgico	X	X			X	X	X	
Consulta primaria a pacientes remotos	X	X		X	X		X	X
Consulta de rutina o de segunda opinión	X	X		X	X	X	X	X
Transmisión de imágenes diagnósticas	X	X		X	X	X	X	X
Control de diagnósticos ampliados	X	X			X			
Manejo de enfermedades crónicas	X	X			X		X	
Transmisión de datos médicos	X	X	X	X			X	X
Salud Pública, medicina preventiva y educación al paciente	X	X		X	X			
Educación y actualización de profesionales de la salud	X	X		X	X			X

Tabla 12.1.: Aplicaciones de la telemedicina según su ámbito de utilización. Basada en [203], p. 75

Equipos Aplicaciones	Videoconferencia	Cámara Digital, Cámara Análoga	Digitalizador RX, Frame Grabber, DICOM	Periféricos de laboratorio	EEG, ECG, ED, Signos Vitales	Dermatoscopio	Oftalmoscopio	Objetivos ORL	Teléfono RTPC o Móvil (Voz)	TV, teleconferencia	Internet
Evaluación inicial del estado de urgencia y transferencia	X	X	X	X	X				X		
Tratamiento médico y post-quirúrgico	X	X		X	X				X		X
Consulta primaria a pacientes remotos	X					X			X		X
Consulta de rutina o de segunda opinión	X			X					X		X
Transmisión de imágenes diagnósticas		X	X			X	X	X			
Control de diagnósticos ampliados	X								X		
Manejo de enfermedades crónicas	X			X	X				X		X
Transmisión de datos médicos				X	X						
Salud Pública, medicina preventiva y educación al paciente	X								X	X	X
Educación y actualización de profesionales de la salud	X									X	X

Tabla 12.2.: Equipos por aplicaciones de telemedicina. Basada en [203], p. 76

12.2. La calidad de servicio (QoS) para las aplicaciones de tiempo real

En la sección anterior se explicaron las características de los servicios de telemedicina de tiempo real y de tiempo diferido. En esta sección se analizarán las características que debe tener la red utilizada por estos servicios para poder ofrecerlos con una calidad adecuada. En primer lugar, hay que tener en cuenta que la calidad de servicio depende de la percepción del usuario y de su expectativa del servicio que se le está ofreciendo. Esta expectativa muchas veces está condicionada por la experiencia previa del usuario. Veamos un par de ejemplos. Una llamada de voz sobre IP que puede parecerle de mala calidad a alguien acostumbrado a usar la telefonía tradicional sobre la red conmutada telefónica, puede parecerle de calidad adecuada a alguien que hasta ese momento no tenía forma de comunicarse. Por otra parte, en la utilización de un teleestetoscopio, la experiencia anterior o la expectativa que tenga quien lo utilice respecto de la calidad del sonido que espera puede determinar que se deba ser más o menos exigente con la calidad de servicio.

Es muy difícil por lo tanto establecer la calidad de servicio adecuada en términos generales, ya que esto depende de parámetros que muchas veces son subjetivos. Si bien existen estandarizaciones de procedimientos para definir la calidad percibida por el usuario, por lo general estos resultan costosos y poco prácticos cuando se trata de aplicaciones en línea o de tiempo real. Por este motivo, habitualmente se trata de mapear la calidad de servicio a un conjunto de parámetros objetivos medibles en la red, como pueden ser el caudal de datos que obtiene la aplicación, el retardo, las pérdidas de paquetes, la variabilidad del retardo (*jitter*), etc. Sin embargo, es difícil realizar el mapeo anterior, y los valores “adecuados” de estos parámetros dependen no sólo del servicio o aplicación sino también de muchos otros factores.

Una primera clasificación de los servicios, a los efectos de definir los parámetros de la red que más impactan en la calidad de servicio, diferencia entre servicios elásticos (que en general se corresponden a los servicios de tiempo diferido vistos en la sección anterior) e inelásticos (donde están incluidos los servicios de telemedicina de tiempo real vistos en la sección anterior). Los servicios o aplicaciones elásticas son aquellas que manejan flujos donde no hay una restricción temporal estricta. Es el ejemplo del correo electrónico, la transferencia de archivos, o aplicaciones de telemedicina como el estudio de un centellograma o un estudio radiológico donde el especialista recibe las imágenes, las diagnostica y envía luego el resultado. Este tipo de servicios pueden “prolongarse” en el tiempo sin que se pierdan totalmente su valor como tal.

Por el contrario, en las aplicaciones inelásticas cada paquete debe llegar a destino y reproducirse antes de un cierto tiempo que es relativamente rígido. En este último caso se tienen por ejemplo las conversaciones telefónicas, las videoconferencias, la teleestetoscopia, la teleecografía u otro tipo de diagnóstico médico donde exista relación entre dos o más agentes que interactúan a través de la red en tiempo real.

Hay que observar que la clasificación anterior no depende del tipo de contenido. Por ejemplo, no depende de que sea audio o video sino del tipo de interactividad requerida

por la aplicación. Con contenidos de tipo audio o video se puede tener una gran gama de requerimientos de calidad diferentes: desde enviar un archivo para que sea escuchado y analizado posteriormente por un especialista, que es una aplicación de tipo elástico, hasta mantener una conversación telefónica o hacer una teleecografía, que son aplicaciones típicamente inelásticas.

En las aplicaciones elásticas el parámetro relevante es el caudal que obtiene la aplicación de la red, y esto está definido por la red y por las características del protocolo TCP, ya que las aplicaciones elásticas se transportan habitualmente sobre este protocolo de capa 4.

En cambio, en las aplicaciones inelásticas impactan también otros parámetros como el retardo y las pérdidas de paquetes. Los parámetros más relevantes dependen del tipo de servicio y también de otras características como los codificadores/decodificadores (códecs) de voz y video que utiliza la aplicación. En aplicaciones interactivas de audio como la telefonía, se recomienda tener retardos menores a 150 ms para mantener una conversación de buena calidad. Sin embargo, en este tipo de aplicaciones las pérdidas de paquetes tienen una menor influencia ya que si se pierde un paquete la aplicación puede, por ejemplo, rellenar ese breve intervalo con el ruido ambiente, y esta degradación no ser muy perceptible para el usuario (obviamente si las pérdidas se mantienen dentro de márgenes razonables, por ejemplo menores al 1 %).

En las aplicaciones inelásticas que involucran video interactivo las pérdidas en general juegan un rol mucho más importante. No es fácil establecer límites generales ya que estos dependen del códec que se utilice. El video consume mucho ancho de banda y por lo tanto se comprime para enviarlo por la red. La forma básica de reducir la tasa de bits que genera un video consiste en utilizar la redundancia que presentan las imágenes y las secuencias. Si bien los codificadores de video pueden utilizar muchas técnicas diferentes de compresión, la base de todas ellas radica en enviar un cuadro entero cada cierto tiempo y luego las diferencias entre el cuadro completo y los siguientes cuadros. La pérdida de paquetes puede impedir entonces la reconstrucción adecuada de la imagen hasta tanto no se reponga un cuadro entero. Por consiguiente, el grado de dependencia de la calidad de servicio de un video interactivo con relación a las pérdidas de paquetes está relacionada con la función específica de compresión y codificación que utilice el códec en cada aplicación.

El *jitter* es también una variable importante porque tanto las muestras de voz como los cuadros de video deben reproducirse en una secuencia de tiempo precisa. Si el retardo es variable, cuando se deba reproducir el siguiente cuadro puede suceder que este aún no hubiere llegado al destinatario. Para solucionar este problema, se utilizan almacenes temporales (*buffers*) en el receptor. Esta técnica, si bien soluciona el problema del *jitter*, tiene una aplicación muy limitada en servicios interactivos ya que aumenta el retardo, que tampoco es deseable.

Por último, es necesario tener en cuenta algunas características particulares de las redes inalámbricas, que son las que habitualmente se utilizan en aplicaciones de telemedicina en zonas rurales. El aire como medio de transmisión es mucho menos confiable que la fibra óptica o el par de cobre. Por lo tanto, la tasa de errores en el aire es mucho mayor que en las redes cableadas. En las redes cableadas las pérdidas de paquetes habitual-

mente se deben al desborde de los almacenes temporales al congestionarse la red. En cambio en las redes inalámbricas hay un fenómeno adicional tan o más importante que el anterior que son las pérdidas por errores de transmisión. Para acotar el efecto de las pérdidas en las redes inalámbricas, habitualmente se introducen mecanismos de detección de errores y retransmisión en la capa 2. Este tipo de mecanismos, si bien mejoran la tasa de errores, agregan un retardo adicional para aquellos paquetes que deben ser retransmitidos y por tanto generan *jitter*. Un efecto similar se tiene en redes de acceso aleatorio como las redes 802.11, ya que en ellas se introducen mecanismos de espera aleatoria para acceder al medio de comunicación, que también agregan retardo y *jitter*.

Por las consideraciones anteriores es sumamente importante analizar las características de la red sobre la que se van a ofrecer servicios de telemedicina, analizar los requerimientos de los usuarios de las aplicaciones específicas a desplegar y luego mapear esos requerimientos a restricciones en los parámetros de la red, como caudal mínimo, retardo, pérdidas y valores de *jitter* máximos admisibles. Una vez realizado este análisis habrá que definir la arquitectura de red a utilizar y en particular qué mecanismos de calidad de servicio se deben implementar para alcanzar los objetivos propuestos. Los mecanismos disponibles dependen de la tecnología específica y se hizo referencia a ellos en los capítulos donde se analizaron las redes 802.11 (Wi-Fi), 802.16 (WiMAX) y satelitales.