

Towards Computer Understanding of Human Interactions

Iain McCowan, Daniel Gatica-Perez, Samy Bengio, Darren Moore, and
Hervé Bourlard

Dalle Molle Institute for Perceptual Artificial Intelligence (IDIAP)
P.O. Box 592, CH-1920, Martigny, Switzerland
{mccowan,gatica,bengio,moore,bourlard}@idiap.ch,
<http://www.idiap.ch/>

Abstract. People meet in order to interact - disseminating information, making decisions, and creating new ideas. Automatic analysis of meetings is therefore important from two points of view: extracting the information they contain, and understanding human interaction processes. Based on this view, this article presents an approach in which relevant information content of a meeting is identified from a variety of audio and visual sensor inputs and statistical models of interacting people. We present a framework for computer observation and understanding of interacting people, and discuss particular tasks within this framework, issues in the meeting context, and particular algorithms that we have adopted. We also comment on current developments and the future challenges in automatic meeting analysis.

1 Introduction

The domain of human-computer interaction aims to help humans interact more naturally with computers. A related emerging domain of research instead views the computer as a tool to assist or understand human interactions : putting computers in the human interaction loop [1]. Humans naturally interact with other humans, communicating and generating valuable information. The most natural interface for entering this information into a computing system would therefore be for the computer to extract it directly from observing the human interactions.

The automatic analysis of human interaction is a rich research area. There is growing interest in the automatic understanding of group behaviour, where the interactions are defined by individuals playing and exchanging both similar and complementary roles (e.g. a handshake, a dancing couple, or a children's game) [2,3,4,5,6]. Most of the previous work has relied on visual information and statistical models, and studied three specific scenarios: surveillance in outdoor scenes [5,6], workplaces [3,4], and indoor group entertainment [2]. Beyond the use of visual information, dialogue modelling [7,8] analyses the structure of interactions in conversations.

While it has only recently become an application domain for computing research, observation of human interactions is not a new field of study - it has been actively researched for over fifty years by a branch of social psychologists [9,10,11]. For example, research has analysed turn-taking patterns in group discussions [12,13,14], giving insight into issues such as interpersonal trust, cognitive load in interactions, and patterns of dominance and influence [11]. Research has also shown that interactions are fundamentally multimodal, with participants coordinating speaking turns using a variety of cues, such as gaze, speech back-channels, changes in posture, etc. [12,13,15]. In general, visual information can help disambiguate audio information [16], and when the modalities are discrepant, participants appear to be more influenced by visual than by audio cues [11,17].

Motivated therefore by a desire to move towards more natural human-machine interfaces, and building upon findings of social psychologists regarding the mechanisms and significance of human interactions, this article presents an observational framework for computer understanding of human interactions, focussing on small group meetings as a particular instance.

Meetings contain many complex interactions between people, and so automatic meeting analysis presents a challenging case study. Speech is the predominant modality for communication in meetings, and speech-based processing techniques, including speech recognition, speaker identification, topic detection, and dialogue modelling, are being actively researched in the meeting context [18,8,19,20]. Visual processing, such as tracking people and their focus of attention, has also been examined in [1,21]. Beyond this work, a place for analysis of text, gestures, and facial expressions, as well as many other audio, visual and multimodal processing tasks can be identified within the meeting scenario. While important advances have been made, to date most approaches to automatic meeting analysis have been limited to the application of known technologies to extract information from individual participants (e.g. speech, gaze, identity, etc). Intuitively, the true information of meetings is created from interactions between participants, and true understanding of meetings can only emerge from considering their group nature.

The remainder of this article is organised as follows. Section 2 describes a multi-sensor meeting room that we have installed to enable our research. A framework for computer understanding of human interactions is outlined in Section 3, along with some specific issues and algorithms related to the meeting context. Finally, some perspective on future directions in automatic meeting analysis is given in Section 4, followed by concluding remarks in Section 5.

2 A Multi-sensor Meeting Room

As mentioned above, interactions between people in meetings are generally multimodal in nature. While the *audio* modality is the most obvious source of information in discussions, studies have shown that significant information is conveyed in the *visual* modality, through expressions, gaze, gestures and posture [12,13,15].

In meetings, the *textual* modality is also important, with presentation slides, whiteboard activity, and shared paper documents providing detailed information.

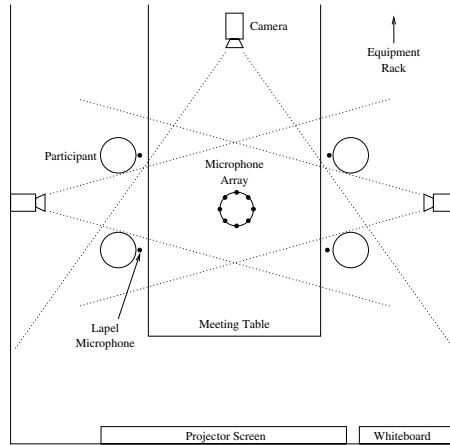


Fig. 1. Meeting recording configuration

To facilitate research into automatic meeting analysis, a meeting room at IDIAP has been equipped with multi-media recording facilities. The audio information is captured using up to 24 microphone channels, including microphone arrays, a binaural manikin, and lapel microphones attached to each meeting participant. Visual information is acquired using three closed circuit television cameras equipped with wide angle lenses. These cameras capture frontal video of meeting participants, as well as wide angle views of the entire meeting room. Textual information is also acquired from presentations and whiteboard usage. Presentation slides are captured from a ceiling-mounted data projector at native VGA resolutions and all whiteboard activity is acquired using transmitting pens and a receiver attached to a standard whiteboard. The acquisition of all modalities is completely synchronised and all data streams are accurately time-stamped.

To date, meeting recording efforts at IDIAP have focussed on the compilation of an audio-visual corpus of approximately sixty, five-minute, four-person scripted meetings. The meeting room configuration used for these recordings is illustrated in Figure 1. Two cameras were used to capture frontal views of the meeting participants (including the table region used for note-taking), while the third camera recorded a view of the whiteboard and presentation screen at the front of the room. An eight-element circular equi-spaced microphone array of 20cm diameter was centrally located on the meeting table.

The resulting five hours of multi-channel, audio-visual meeting data is available for public distribution through a MultiModal Media file server at `mmm.idiap.ch`. A new round of meeting collection will be launched in the near future, and will utilise the recently added slide and whiteboard capture capabilities.

3 Multimodal Processing

We propose a framework for computer understanding of human interactions that involves the following basic steps in a processing loop :

1. locate and track participants
2. for each located participant
 - a) enhance their audio and visual streams
 - b) identify them
 - c) recognise their individual actions
3. recognise group actions

The first step is necessary to determine the number and location of participants. For each person present, we then extract a dedicated enhanced audio and visual stream by focussing on their tracked location. Audio-visual (speech and face) speaker identification techniques can then be applied to determine who the participant is. Individual actions, such as speech activity, gestures or speech words may also be measured or recognised from the audio and visual streams. The ultimate goal of this analysis is then to be able to recognise actions belonging to the group as a whole, by modelling the interactions of the individuals.

Specific issues and algorithms for implementing a number of these steps for the case of meeting analysis are presented in the following sub-sections. A primary focus of our research is the multimodal nature of human interactions in meetings, and this is reflected in the choice of tasks we have included. Naturally, there are many other processing tasks involved in understanding meetings, such as speech recognition and dialogue modelling, that are not covered here.

3.1 Audio-visual Speaker Tracking

The problem in the global view. Locating and tracking speakers represents an important first step towards automatic understanding of human interactions. As mentioned previously, speaker turn patterns convey a rich amount of information about the behaviour of a group and its individual members [10,13]. Furthermore, experimental evidence has highlighted the role that non-verbal behaviour (gaze, facial expressions, and body postures) plays in interactions [13]. Recognising such rich multimodal behaviour first requires reliable localisation and tracking of people.

Challenges in the meeting context. The separate use of audio and video as cues for tracking are classic problems in signal processing and computer vision. However, sound and visual information are jointly generated when people speak, and provide complementary advantages. While initialisation and recovery from failures can be addressed with audio, precise object localisation is better suited to visual processing.

Long-term, reliable tracking of multiple people in meetings is challenging. Meeting rooms pose a number of issues for audio processing, such as reverberation and multiple concurrent speakers, as well as for visual processing, including clutter and variations of illumination. However, the main challenge arises from the behaviour of multiple participants resulting in changes of appearance and pose for each person, and considerable (self)-occlusion. At the same time, meetings in a multi-sensor room present some advantages that ease the location and tracking tasks. Actions usually unfold in specific areas (meeting table, whiteboard, and projector screen), which constrains the group dynamics in the physical space. In addition, the availability of multiple cameras with overlapping fields of view can be exploited to build more reliable person models, and deal with the occlusion problems.

Our approach. We have developed a principled method for speaker tracking, fusing information coming from multiple microphones and uncalibrated cameras [22], based on *Sequential Monte Carlo* (SMC) methods, also known as *particle filters* (PFs) [23]. For a state-space model, a PF recursively approximates the conditional distribution of states given observations using a dynamical model and random sampling by (i) generating candidate configurations from the dynamics (*prediction*), and (ii) measuring their likelihood (*updating*), in a process that amounts to random search in a configuration space. Data fusion can be introduced in both stages of the PF algorithm.

Our work is guided by inherent features of AV data. First, audio is a strong cue to model discontinuities that clearly violate usual assumptions in dynamics (including speaker turns across cameras), and (re)initialisation. Its use for prediction thus brings benefits to modelling real situations. Second, audio can be inaccurate at times, but provides a good initial localisation guess that can be enhanced by visual information. Third, although audio might be imprecise, and visual calibration can be erroneous due to distortion in wide-angle cameras, the joint occurrence of AV information in the constrained physical space in meetings tends to be more consistent, and can be learned from data.

Our methodology exploits the complementary features of the AV modalities. In the first place, we use a 2-D approach in which human heads are visually represented by their silhouette in the image plane, and modelled as elements of a *shape-space*, allowing for the description of a head template and a set of valid geometric transformations (motion). In the second place, we employ a *mixed-state space*, where in addition to the continuous subspace that represents head motion, we include a discrete component that indicates the specific camera plane in which a speaker is present. This formulation helps define a generative model for camera

switching. In the third place, we asymmetrically handle audio and video in the PF formulation. Audio localisation information in 3-D space is first estimated by an algorithm that reliably detects speaker changes with low latency, while maintaining good estimation accuracy. Audio and skin-color blob information are then used for prediction, and introduced in the PF via *importance sampling*, a technique which guides the search process of the PF towards regions of the state space likely to contain the true configuration (a speaker). Additionally, audio, color, and shape information are jointly used to compute the likelihood of candidate configurations. Finally, we use an AV calibration procedure to relate audio estimates in 3-D and visual information in 2-D. The procedure uses easily generated training data, and does not require precise geometric calibration of cameras and microphones [22].



Fig. 2. Tracking speakers in the meeting room. Frames 100, 1100, 1900, and 2700.

The result is a method that can initialise and track a moving speaker, and switch between multiple people across cameras with low delay, while tolerating visual clutter. An example for the setup of Figure 1 is shown in Figure 2. For a two-minute sequence, the system tracked the current speaker in the correct camera in approximately 88% of the frames, while keeping the localisation error in the corresponding image plane within a few pixels. Other AV tracking examples for single- and multi-camera set-ups can be found at www.idiap.ch/~gatica.

Open problems. Although the current methodology is useful in its current form, there is much room for improvement. In the following we identify three specific lines of research. We are currently generalising our formulation to a multiple-object AV tracker, which involves the integration of person-dependent appearance models, and the consistent labelling of tracked objects along time and across cameras. Multi-object tracking significantly increases the dimensionality of the state space, which calls for efficient inference mechanisms in the resulting statistical model. Another line of research is the integration of more robust person models. The third line of research is the joint formulation of tracking and recognition. Specifically, we are building head trackers that simultaneously estimate head orientation (a simple form of recognition), which is in turn a strong cue for detection of focus of attention, and useful for higher-level recognisers.

3.2 Speech Segmentation and Enhancement Using Microphone Arrays

The problem in the global view. Having located and tracked each person, it is next necessary to acquire an enhanced dedicated audio channel of their speech. Speech is the predominant communication modality, and thus a rich source of information, in many human interactions.

Most state-of-the-art speech and speaker recognition systems rely on close-talking head-set microphones for speech acquisition, as they naturally provide a higher signal-to-noise ratio (SNR) than single distant microphones. This mode of acquisition may be acceptable for applications such as dictation, however as technology heads towards more pervasive applications, less constraining solutions are required. *Microphone arrays* present a promising alternative to close-talking microphones, as they allow for signal-independent enhancement, localisation and tracking of speakers, and non-intrusive hands-free operation. For these reasons, microphone arrays are being increasingly used for speech acquisition in such applications [24,25].

Challenges in the meeting context. Meetings present a number of interesting challenges for microphone array research. A primary issue is the design of the *array geometry* : how many microphones should be used, and where should they be placed in the room? Naturally a geometry giving high spatial resolution uniformly across a room is desirable for best performance and lowest constraint on the users, however this requires prohibitively large numbers of microphones, and complex installation [26]. For these reasons, more practical solutions with smaller numbers of microphones need to be researched to address computational and economical considerations.

A second challenge in the meeting context is the natural occurrence of overlapping speech. In [27] it was identified that around 10-15% of words, or 50% of speech segments, in a meeting contain a degree of overlapping speech. These overlapped segments are problematic for speaker segmentation, and speech and

speaker recognition. For instance, an absolute increase in word error rate of between 15-30% has been observed on overlap speech segments using close-talking microphones [27,8].

Our approach. While it is clear that a large microphone array with many elements would give the best spatial selectivity for localisation and enhancement, for microphone arrays to be employed in practical applications, hardware cost (microphones, processing and memory requirements) must be reduced. For this reason, we focus on the use of small microphone arrays, which can be a viable solution when assumptions can be made about the absolute and relative locations of participants.

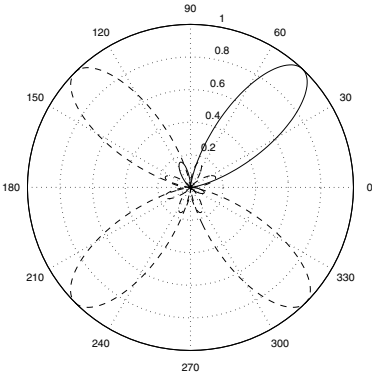


Fig. 3. Microphone array directivity patterns at 1000 Hz (speaker 1 direction in bold)

As shown in Figure 1, the particular array geometry we have chosen is an 8-element circular array (of radius 10cm) placed at the centre of the meeting table. This geometry and placement was selected based on the assumption that a meeting generally consists of small groups of people seated and talking face to face in well-defined regions. Each array is designed to cater for a small group of up to 4 people. In larger meetings, multiple (potentially interacting) small array modules are positioned along the table, where each module is responsible for the people in its local region. The circular geometry was selected as it gives uniform spatial selectivity between people sitting around it, leading to good general performance in separating overlapping speech. This is important for meetings where background noise is generally low, and so overlapping speech is the primary noise source. To illustrate, Figure 3 shows the theoretical *directivity pattern* (array gain as a function of direction) for the array at 1000 Hz for 4 speakers separated by 90 degrees. Having the array on the table also means it is placed in close proximity to participants, leading to naturally high signal levels compared to background noise caused by distant sources.

Given accurate tracking of the speaker locations in the room, the next task is to determine segments of continuous speech from a given speaker location. *Speaker segmentation* in meetings is problematic for traditional techniques based on simple energy or spectral features, as a significant amount of cross-talk from other speakers exists even on close-talking microphones [28,29]. In [30,31] we presented a *location-based segmentation* technique that is capable of providing a smooth speech/silence segmentation for a given room location. As it is based on speech location features from the microphone array, rather than standard spectral features, this location-based segmentation has the important benefit of being able to accurately handle multiple concurrent speakers (identifying which locations are active at any given time). As an example, our current system is capable of segmenting four person meetings, including overlapping speech segments, with over 95% frame accuracy [31].

Once the location of the speakers is known along with their speech activity segmentation, we can then apply microphone array *beamforming* techniques to enhance their speech, attenuating background noise and conflicting speech sources. Beamforming consists of filtering and combining the individual microphone signals in such a way as to enhance signals coming from a particular location. For beamforming filters, we adopt standard *superdirective* filters, which are calculated to maximise the array gain for the desired direction [32]. In addition, we apply a *Wiener post-filter* to the beamformer output to further reduce the broadband noise energy. The post-filter is estimated from the auto- and cross-spectral densities of the microphone array inputs, and is formulated assuming a diffuse background noise field [33]. This post-filter leads to significant improvements in terms of SNR and speech recognition performance in office background noise [33], and has also been shown to out-perform lapel microphones for a small vocabulary recognition task in significant levels of overlapping speech [34].

Open problems. While microphone array speech processing techniques are already relatively mature, a number of open issues remain in this context. As mentioned briefly, larger meetings could be catered for by a series of small microphone array modules working together. A current focus of our research is thus to propose algorithms for these interactions between modules. Further research is also focussing on issues related to the real-time implementation of multiple concurrent beamformers.

3.3 Audio-visual Person Identification

The Problem in the Global View. Identifying participants is important for understanding human interactions. When prior knowledge about the participants is available (such as their preferred way of communicating, topics of interests, levels of language, relative hierarchical levels in a given context, etc), knowing the participants' identities would imply knowing this prior information, which could in turn be used to better tune the algorithms used to analyse the interaction. Fortunately, *biometric authentication* [35], which is the general problem of authenticating or identifying a person using his or her behavioural and

physiological characteristics such as the face or the voice, is a growing research domain which has already shown useful results, especially when using more than one of these characteristics, as we propose to do here.

Challenges in the Meeting Context. In order to perform AV identification during a meeting, we need to extract reliably the basic modalities. For the face, we require a face localisation algorithm that is robust to the kind of images available from a video stream (relatively low-quality and low-resolution), robust to the participants' varying head poses, and able to cope with more than one face per image. This could be done using our AV tracking system described in Section 3.1. For the voice, taking into account that several microphones are available in the meeting room, the first challenge is to separate all audio sources and attribute each speech segment to its corresponding participant. Again, this could be done using our speaker segmentation and enhancement techniques, described in Section 3.2. Afterward, classical face and speaker verification algorithms could be applied, followed by a fusion step, which provides robustness to the failure of one or the other modality. Finally, an identification procedure could be applied.

Our Approach. Our identification system is based on an AV biometric verification system. Assuming that we are able to obtain reliable speech segments and localised faces from the meeting raw data, we can then apply our state-of-the-art verification system, which is based on a *speaker verification* system, a *face verification* system, and a *fusion* module.

Our speaker verification system first starts by extracting useful features from the raw speech data: we extract 16 Linear Predictive Cepstral Coefficient (LPCC) features every 10 ms, as well as their first temporal derivative. Then, a silence detector based on an unsupervised 2-Gaussian system is used to remove all silence frames. Finally, the verification system itself is based on the modelling of one Gaussian Mixture Model (GMM) for each individual, adapted using *Maximum A Posteriori* (MAP) techniques from a *World Model* trained by *Expectation-Maximisation* on a large set of prior data. The score for a given access is obtained as the logarithm of the ratio between the likelihood of the data given the individual model and the likelihood given the world model. This system obtains state-of-the-art performance on several benchmark verification databases [36].

Our face verification system is based on a non-holistic view: instead of extracting features from a full face image that are then handled by a classifier, as it is often done [37], we extract Discrete Cosine Transform (DCT) -based features from overlapping blocks of the image (square patches that span the whole face image), using the DCTmod2 technique [38], which has shown state-of-the-art performance in several cases. Finally, we model the obtained feature vectors using GMMs, similarly to the speaker verification system, hence using the same scoring technique.

Our fusion algorithm is based on Multi-layer Perceptrons (experiments with Support Vector Machines give similar performances). The fusion model takes as

input the log likelihood scores coming from both the face and the speaker verification systems, and combines them non-linearly in order to obtain a unified and more robust overall score. Optionally, confidence values could also be computed on both the voice and face scores, which then enhance the quality of the fusion model [39].

Finally, in order to identify the correct individual, the whole verification system is run over all previously stored individual models, and the model corresponding to the highest obtained score over a pre-defined threshold (in order to account for unknown individuals) identifies the target individual.

Open Problems. Assuming that speaker segmentation and face tracking have given perfect segmentation, for a given meeting, we will have potentially several minutes of speech and face data per individual. In general, a classical verification system only requires a few face images and less than one minute of speech data to attain acceptable performance. However, the environment is unconstrained, the meeting data may be noisy for different reasons - the individual may not always look at the camera and speak loudly and intelligibly. In this case, rather than using all available data to identify a person, a better solution could be to be more strict on the selection of faces and speaker segments in order to keep only the best *candidates* for identification. Hence, we should try to remove highly noisy or overlapping speech segments, badly tracked face images and faces that are not in a good frontal pose and good lighting condition.

3.4 Group Action Recognition

The problem in the global view. The ultimate goal of automatic analysis of human interactions is to recognise the group actions. As discussed previously, the true information of meetings is created from interactions between participants playing and exchanging roles. In this view, an important goal of automatic meeting analysis is the segmentation of meetings into high-level agenda items which reflect the action of the group as a whole, rather than just the behaviour of individuals (e.g. discussions and presentations, or even higher level notions, like planning, negotiating, and making decisions).

Challenges in the meeting context. Recognition of group actions in meetings entails several important problems for which no satisfactory solutions currently exist. These include (1) devising tractable multi-stream sequence models, where each stream could arise from either a modality (AV) or a participant; (2) modelling asynchronicity between participants' behaviour; (3) extracting features for recognition that are robust to variations in human characteristics and behaviour; (4) designing sequence models that can integrate language features (e.g. keywords or dialog acts) with non-verbal features (e.g. emotion as captured from audio and video); and (5) developing models for recognition of actions that are part of a hierarchy.

One potentially simplifying advantage to recognise group actions in meetings is that participants usually have some influence on each other's behaviour. For example, a dominant speaker grabbing the floor often makes the other participants go silent, and a presentation will draw most participants' attention in the same direction. The recognition of some group actions can be therefore benefit from the occurrence of these multiple similar individual behaviours.

Our approach. We have addressed meeting group action recognition as the recognition of a continuous, non-overlapping, sequence of lexical entries, analogous to observational approaches in social psychology for analysis of group interaction [10], and to speech or continuous gesture recognition [40,41]. Continuous recognition generates action-based meeting segmentations that can be directly used for browsing. Furthermore, the definition of multiple lexica would provide alternative semantic views of a meeting. Note that in reality, most group actions are characterised by soft (natural) transitions, and specifying their boundaries beyond a certain level of precision has little meaning.

In particular, we have modelled meeting actions based on a set of multimodal turn-taking events. Speaking turns are mainly characterised by audio information, but significant information is also present in non-verbal cues like gaze and posture changes [13], which can also help disambiguate audio information [16]. The specific actions include monologues (one participant speaks continuously without interruption), discussions (all participants engage in a discussion), presentations (one participant at front of room makes a presentation using the projector screen), white-boards (one participant at front of room talks and uses the white-board), and group note-taking (all participants write notes).

To investigate the multimodal and group natures of the actions, we used a variety of Hidden Markov Models (HMMs) [40] to combine the streams of information (with streams representing modalities or people) in different ways. The models include early integration HMMs, multi-stream HMMs [42], and asynchronous HMMs [43]. Furthermore, the individual behaviour of participants was monitored using features from both the audio and visual modalities (including speech activity, pitch, energy, speaking rate, and head and hand location and motion features).

A detailed account of our experiments and results can be found in [44]. For experiments, we used the meeting corpus described in Section 2. Meetings followed a loose script to ensure an adequate amount of examples of all actions, and to facilitate annotation for training and testing, but otherwise the individual and group behaviour is natural. In summary, the best action error rate (equivalent to the word error rate in speech recognition) that we obtained on an independent test set was 5.5% using a two-stream HMM, where one stream modelled audio features, and the other modelled video features coming from all participants. Several other models (audio-only, AV early-integration, and AV asynchronous) produced competitive results. As expected, for this set of actions audio was the main source of information for reliable recognition, while video mostly helped in reducing monologue and discussion errors.

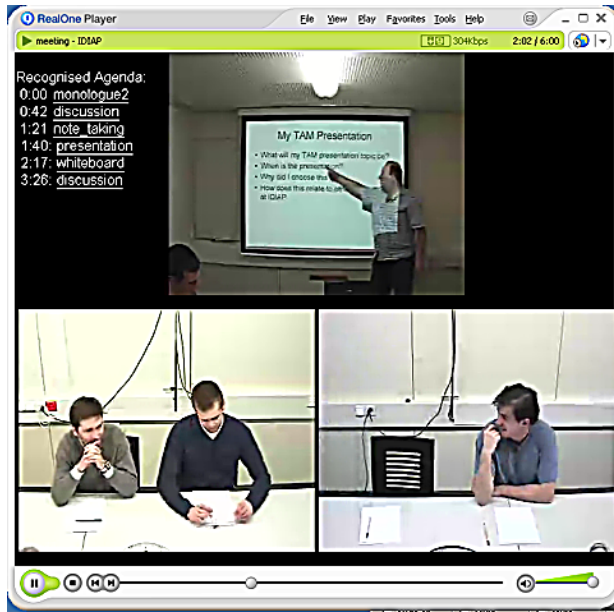


Fig. 4. Simple meeting browser interface, showing recognised meeting actions.

An example of the application of the action recognition results for meeting browsing is shown in Figure 4.

Open problems. The experience gained from our results confirms the importance of modelling the interactions between individuals, as well as the advantage of a multimodal approach for recognition. We believe there is much scope for work towards the recognition of different sets of high-level meeting actions, including other multimodal turn-taking events, actions based on participants' mood or level of interest, and multimodal actions motivated by traditional dialogue acts. To achieve this goal, ongoing and future work will investigate richer feature sets, and appropriate models for the interactions of participants. Another task will be to incorporate prior information in the recognition system, based on the participant identities and models of their personal behaviour. We also plan to collect a larger meeting corpus, and work on the development of more flexible assessment methodologies.

4 Future Directions

From the framework outlined in the beginning of Section 3, while much room clearly remains for new techniques and improvements on existing ones, we can see that steps 1-2(c) are reasonably well understood by the state-of-the-art. In

contrast, we are far from making similar claims regarding step 3, recognition of group actions.

The first major goal in computer understanding of group actions, is to clearly identify lexica of such actions that may be recognised. A simple lexicon based on multimodal turn-taking events was discussed in Section 3.4, however there is a need to progress towards recognition of higher level concepts, such as decisions, planning, and disagreements. In this regard, the social psychology literature represents an important source of information for studies on the tasks and processes that arise from human interactions, as was discussed in [44].

Having identified relevant group actions, a further research task is then to select appropriate features for these actions to be recognised. At this moment, features are intuitively selected by hand, which has obvious limitations. Approaches for feature selection could arise from two areas. The first one is human. We require a deeper understanding of human behaviour. Existing work in psychology could provide cues for feature selection towards, for example, multimodal recognition of emotion [45]. The second one is computational. Developments in machine learning applied to problems in vision and signal processing point to various directions [46].

Finally, to recognise the group actions, there is a need to propose models capable of representing the interactions between individuals in a group (see e.g. [47,5,44]). Some particular issues are the need to model multiple data streams, asynchronicity between streams, hierarchies of data and events, as well as features of different nature (e.g. discrete or continuous).

5 Conclusion

This article has discussed a framework for computer understanding of human interactions. A variety of multimodal sensors are used to observe a group and extract useful information from their interactions. By processing the sensor inputs, participants are located, tracked, and identified, and their individual actions recognised. Finally, the actions of the group as a whole may be recognised by modelling the interactions of the individuals.

While initial work in this direction has already shown promising progress and yielded useful results, it is clear that many research challenges remain if we are to advance towards true computer understanding of human interactions.

Acknowledgements. The authors would like to acknowledge our colleagues at IDIAP involved in the research described in this article, in particular Guillaume Lathoud, Johnny Mari  thoz, Sebastien Marcel, Conrad Sanderson, Olivier Masson, Pierre Wellner, Jean Marc Odobez and Florent Monay.

This work was supported by the Swiss National Science Foundation through the National Centre of Competence in Research (NCCR) on “Interactive Multimodal Information Management (IM2)”. The work was also funded by the European project “M4: MultiModal Meeting Manager”, through the Swiss Federal Office for Education and Science (OFES).

References

1. A. Waibel, T. Schultz, M. Bett, R. Malkin, I. Rogina, R. Stiefelhagen, and J. Yang, "SMaRT:the Smart Meeting Room Task at ISL," in *Proc. IEEE ICASSP 2003*, 2003.
2. A. Bobick, S. Intille, J. Davis, F. Baird, C. Pinhanez, L. Campbell, Y. Ivanov, A. Schutte, and A. Wilson, "The KidsRoom: A Perceptually-Based Interactive and Immersive Story Environment," *PRESENCE: Teleoperators and Virtual Environments*, vol. 8, August 1999.
3. N. Johnson, A. Galata, and D. Hogg, "The acquisition and use of interaction behaviour models," in *Proc. IEEE Int. Conference on Computer Vision and Pattern Recognition*, June 1998.
4. T. Jebara and A. Pentland, "Action reaction learning: Automatic visual analysis and synthesis of interactive behaviour," in *Proc. International Conference on Vision Systems*, January 1999.
5. N. Oliver, B. Rosario, and A. Pentland, "A bayesian computer vision system for modeling human interactions," *IEEE Transactions on Pattern Analysis and Machine Intelligence*, vol. 22, August 2000.
6. S. Hongeng and R. Nevatia, "Multi-agent event recognition," in *Proc. IEEE Int. Conference on Computer Vision*, (Vancouver), July 2001.
7. J. Carletta, A. Isard, S. Isard, J. Kowtko, G. Doherty-Sneddon, and A. Anderson, "The coding of dialogue structure in a corpus," in *Proceedings of the Twente Workshop on Language Technology: Corpus-based approaches to dialogue modelling* (J. Andernach, S. van de Burgt, and G. van der Hoeven, eds.), Universiteit Twente, 1995.
8. N. Morgan, D. Baron, J. Edwards, D. Ellis, D. Gelbart, A. Janin, T. Pfau, E. Shriberg, and A. Stolcke, "The meeting project at ICSI," in *Proc. of the Human Language Technology Conference*, (San Diego, CA), March 2001.
9. R. F. Bales, *Interaction Process Analysis: A method for the study of small groups*. Addison-Wesley, 1951.
10. J. E. McGrath, *Groups: Interaction and Performance*. Prentice-Hall, 1984.
11. J. McGrath and D. Kravitz, "Group research," *Annual Review of Psychology*, vol. 33, pp. 195–230, 1982.
12. E. Padilha and J. C. Carletta, "A simulation of small group discussion," in *EDIALOG*, 2002.
13. K. C. H. Parker, "Speaking turns in small group interaction: A context-sensitive event sequence model," *Journal of Personality and Social Psychology*, vol. 54, no. 6, pp. 965–971, 1988.
14. N. Fay, S. Garrod, and J. Carletta, "Group discussion as interactive dialogue or serial monologue: The influence of group size," *Psychological Science*, vol. 11, no. 6, pp. 487–492, 2000.
15. D. Novick, B. Hansen, and K. Ward, "Coordinating turn-taking with gaze," in *Proceedings of the 1996 International Conference on Spoken Language Processing (ICSLP-96)*, 1996.
16. R. Krauss, C. Garlock, P. Bricker, and L. McMahon, "The role of audible and visible back-channel responses in interpersonal communication," *Journal of Personality and Social Psychology*, vol. 35, no. 7, pp. 523–529, 1977.
17. B. DePaulo, R. Rosenthal, R. Eisenstat, P. Rogers, and S. Finkelstein, "Decoding discrepant nonverbal cues," *Journal of Personality and Social Psychology*, vol. 36, no. 3, pp. 313–323, 1978.

18. F. Kubala, "Rough'n'ready: a meeting recorder and browser," *ACM Computing Surveys*, vol. 31, 1999.
19. A. Waibel, M. Bett, F. Metze, K. Ries, T. Schaaf, T. Schultz, H. Soltan, H. Yu, and K. Zechner, "Advances in automatic meeting record creation and access," in *Proc. IEEE ICASSP*, (Salt Lake City, UT), May 2001.
20. S. Renals and D. Ellis, "Audio information access from meeting rooms," in *Proc. IEEE ICASSP 2003*, 2003.
21. R. Cutler, Y. Rui, A. Gupta, J. Cadiz, I. Tashev, L. He, A. Colburn, Z. Zhang, Z. Liu, and S. Silverberg, "Distributed meetings: A meeting capture and broadcasting system," in *Proc. ACM Multimedia Conference*, 2002.
22. D. Gatica-Perez, G. Lathoud, I. McCowan, and J.-M. Odobez, "A mixed-state i-particle filter for multi-camera speaker tracking," in *Proceedings of WOMTEC*, September 2003.
23. A. Doucet, N. de Freitas, and N. Gordon, *Sequential Monte Carlo Methods in Practice*. Springer-Verlag, 2001.
24. R. Cutler, "The distributed meetings system," in *Proceedings of IEEE ICASSP 2003*, 2003.
25. V. Stanford, J. Garofolo, , and M. Michel, "The nist smart space and meeting room projects: Signals, acquisition, annotation, and metrics," in *Proceedings of IEEE ICASSP 2003*, 2003.
26. H. Silverman, W. Patterson, J. Flanagan, and D. Rabinkin, "A digital processing system for source location and sound capture by large microphone arrays," in *Proceedings of ICASSP 97*, April 1997.
27. E. Shriberg, A. Stolcke, and D. Baron, "Observations on overlap: findings and implications for automatic processing of multi-party conversation," in *Proceedings of Eurospeech 2001*, vol. 2, pp. 1359–1362, 2001.
28. T. Pfau, D. Ellis, and A. Stolcke, "Multispeaker speech activity detection for the ICSI meeting recorder," in *Proceedings of ASRU-01*, 2001.
29. T. Kemp, M. Schmidt, M. Westphal, and A. Waibel, "Strategies for automatic segmentation of audio data," in *Proceedings of ICASSP-2000*, 2000.
30. G. Lathoud and I. McCowan, "Location based speaker segmentation," in *Proceedings of the International Conference on Acoustics, Speech and Signal Processing*, April 2003.
31. G. Lathoud, I. McCowan, and D. Moore, "Segmenting multiple concurrent speakers using microphone arrays," in *Proceedings of Eurospeech 2003*, September 2003.
32. J. Bitzer and K. U. Simmer, "Superdirective microphone arrays," in *Microphone Arrays* (M. Brandstein and D. Ward, eds.), ch. 2, pp. 19–38, Springer, 2001.
33. I. McCowan and H. Bourlard, "Microphone array post-filter based on noise field coherence," *To appear in IEEE Transactions on Speech and Audio Processing*, November 2003.
34. D. Moore and I. McCowan, "Microphone array speech recognition: Experiments on overlapping speech in meetings," in *Proceedings of the International Conference on Acoustics, Speech and Signal Processing*, April 2003.
35. A. Jain, R. Bolle, and S. Pankanti, *Biometrics: Person Identification in Networked Society*. Kluwer Publications, 1999.
36. J. Mariéthoz and S. Bengio, "A comparative study of adaptation methods for speaker verification," in *Proceedings of the International Conference on Spoken Language Processing, ICSLP*, 2002.
37. S. Marcel and S. Bengio, "Improving face verification using skin color information," in *Proceedings of the 16th International Conference on Pattern Recognition, ICPR*, IEEE Computer Society Press, 2002.

38. C. Sanderson and K. Paliwal, "Polynomial Features for Robust Face Authentication," *Proceedings of International Conference on Image Processing*, vol. 3, pp. 997–1000, 2002.
39. S. Bengio, C. Marcel, S. Marcel, and J. Mariéthoz, "Confidence measures for multimodal identity verification," *Information Fusion*, vol. 3, no. 4, pp. 267–276, 2002.
40. L. R. Rabiner and B.-H. Juang, *Fundamentals of Speech Recognition*. Prentice-Hall, 1993.
41. T. Starner and A. Pentland, "Visual recognition of american sign language using HMMs," in *Proc. Int. Work. on Auto. Face and Gesture Recognition*, (Zurich), 1995.
42. S. Dupont and J. Luettin, "Audio-visual speech modeling for continuous speech recognition," *IEEE Transactions on Multimedia*, vol. 2, pp. 141–151, September 2000.
43. S. Bengio, "An asynchronous hidden markov model for audio-visual speech recognition," in *Advances in Neural Information Processing Systems, NIPS 15* (S. Becker, S. Thrun, and K. Obermayer, eds.), MIT Press, 2003.
44. I. McCowan, D. Gatica-Perez, S. Bengio, and G. Lathoud, "Automatic analysis of multimodal group actions in meetings," Tech. Rep. RR 03–27, IDIAP, 2003.
45. B. De Gelder and J. Vroomen, "The perception of emotions by ear and by eye," *Cognition and Emotion*, vol. 14, pp. 289–311, 2002.
46. P. Viola and M. Jones, "Rapid object detection using a boosted cascade of simple features," in *Proc. IEEE Int. Conf. on Computer Vision (CVPR)*, (Kawaii), Dec. 2001.
47. S. Basu, T. Choudhury, B. Clarkson, and A. Pentland, "Learning human interactions with the influence model," Tech. Rep. 539, MIT Media Laboratory, June 2001.